

Cena 13,00 zł
(VAT 8%)

Indeks 381306
e-ISSN 2543-8476
PL ISSN 0043-518X

WIADOMOŚCI STATYSTYCZNE

THE POLISH STATISTICIAN

MARZEC / MARCH
ROK / VOLUME 66

2021 | 3

GŁÓWNY URZĄD STATYSTYCZNY
STATISTICS POLAND

POLSKIE TOWARZYSTWO STATYSTYCZNE
POLISH STATISTICAL ASSOCIATION



WIADOMOŚCI STATYSTYCZNE

THE POLISH STATISTICIAN

MARZEC / MARCH
ROK / VOLUME 66

2021 | 3 (718)

RADA NAUKOWA / SCIENTIFIC COUNCIL

dr Dominik Rozkrut (przewodniczący/chairman) – Uniwersytet Szczeciński; Prof. Anthony Arundel – Maastricht University; Eric Bartelsman, PhD, Assoc. Prof. – Vrije Universiteit Amsterdam; prof. dr hab. Czesław Domański – Uniwersytet Łódzki; prof. dr hab. Elżbieta Gołata – Uniwersytet Ekonomiczny w Poznaniu; Semen Matkovskiy, PhD, Assoc. Prof. – Ivan Franko National University of Lviv; prof. dr hab. Włodzimierz Okrasa – Uniwersytet Kardynała Stefana Wyszyńskiego w Warszawie; prof. dr hab. Józef Oleński – Polskie Towarzystwo Statystyczne; prof. dr hab. Tomasz Panek – Szkoła Główna Handlowa w Warszawie; Juan Manuel Rodríguez Poo, PhD, Assoc. Prof. – University of Cantabria; Iveta Stankovičová, BEng, PhD, Assoc. Prof. – Comenius University in Bratislava; prof. dr hab. Marek Walesiak – Uniwersytet Ekonomiczny we Wrocławiu; prof. dr hab. Józef Zegar – Instytut Ekonomiki Rolnictwa i Gospodarki Żywnościowej – Państwowy Instytut Badawczy

sekretarz/secretary: Paulina Kucharska-Singh

KOLEGIUM REDAKCYJNE / EDITORIAL BOARD

Tudorel Andrei, PhD, Assoc. Prof. – Bucharest Academy of Economic Studies; mgr Renata Bielak – Główny Urząd Statystyczny; dr Marek Cierpiał-Wolan – Uniwersytet Rzeszowski; dr hab. Grażyna Dehnel, prof. UEP – Uniwersytet Ekonomiczny w Poznaniu; dr Jacek Kowalewski – Uniwersytet Ekonomiczny w Poznaniu; dr Jan Kubacki – Urząd Statystyczny w Łodzi; dr Grażyna Marciniak; dr hab. Andrzej Młodak, prof. AK – Akademia Kaliska im. Prezydenta Stanisława Wojciechowskiego; dr hab. Mateusz Pipień, prof. UEK – Uniwersytet Ekonomiczny w Krakowie; Marek Rojciček, BEng, PhD – University of Economics, Prague; Anna Shostya, PhD, Assoc. Prof. – Pace University in New York; dr hab. Małgorzata Tarczyńska-Łuniewska, prof. US – Uniwersytet Szczeciński; dr Wioletta Wrzaszcz – Instytut Ekonomiki Rolnictwa i Gospodarki Żywnościowej – Państwowy Instytut Badawczy; dr inż. Agnieszka Zgierska – Główny Urząd Statystyczny

ZESPÓŁ REDAKCYJNY / EDITORIAL STAFF

redaktor naczelny / editor-in-chief: Marek Cierpiał-Wolan

zastępca redaktora naczelnego / deputy editor-in-chief: Andrzej Młodak

redaktorzy tematyczni / thematic editors: Jan Kubacki, Małgorzata Tarczyńska-Łuniewska, Agnieszka Zgierska

redaktor merytoryczny / substantive editor: Wioletta Wrzaszcz

sekretarz/secretary: Małgorzata Zygmunt

ADRES REDAKCJI / EDITORIAL OFFICE ADDRESS

Główny Urząd Statystyczny / Statistics Poland, al. Niepodległości 208, 00-925 Warszawa
tel./phone +48 22 608 32 25, e-mail: redakcja.ws@stat.gov.pl

Redakcja językowa: Wydział Czasopism Naukowych, Główny Urząd Statystyczny
Language editing: Scientific Journals Division, Statistics Poland

Redakcja techniczna, skład i łamanie, wykresy, korekta: Zakład Wydawnictw Statystycznych – zespół pod kierunkiem Wojciecha Szuchty

Technical editing, typesetting, figures, proof-reading: Statistical Publishing Establishment – team supervised by Wojciech Szuchta



Zakład Wydawnictw
Statystycznych

Druk i oprawa / Printed and bound:

Zakład Wydawnictw Statystycznych / Statistical Publishing Establishment
al. Niepodległości 208, 00-925 Warszawa, zws.stat.gov.pl

Wersja elektroniczna, stanowiąca wersję pierwotną czasopisma, jest dostępna na ws.stat.gov.pl
The original version of the journal is the electronic issue, available at ws.stat.gov.pl

© Copyright by Główny Urząd Statystyczny

Indeks 381306

Informacje w sprawie sprzedaży czasopisma / Sales of the journal:

Zakład Wydawnictw Statystycznych / Statistical Publishing Establishment

tel./phone +48 22 608 32 10, +48 22 608 38 10

Prenumerata jest prowadzona przez / Subscription is available at RUCH S.A.

Zamówienia na prenumeratę można składać na stronie / Subscriptions can be ordered at
www.prenumerata.ruch.com.pl

SPIS TREŚCI

CONTENTS

Od redakcji	4
From the editorial team	
Studia metodologiczne	
Methodological studies	
Mirosław Szreder	
Błędy nielosowe i ich znaczenie w testowaniu hipotez	7
Non-random errors and their importance in testing of hypotheses	
Statystyka w praktyce	
Statistics in practice	
Dorota Witkowska, Aleksandra Matuszewska-Janica	
Czynniki determinujące dysproporcje w wynagrodzeniach kobiet i mężczyzn w krajach Unii Europejskiej	22
Factors determining disproportions in men and women's wages in the European Union countries	
Studia interdyscyplinarne. Wyzwania badawcze	
Interdisciplinary studies. Research challenges	
Mariusz Malinowski	
Wykorzystanie Google Trends do modelowania stopy bezrobocia rejestrowanego w Polsce	45
Use of Google Trends in modelling registered unemployment rate in Poland	
In memoriam	
Dominik Rozkrut, Bożena Łazowska, Czesław Domański, Włodzimierz Okrasa, Irena E. Kotowska, Janusz Czapiński, Waldemar Tarczyński, Zofia Kozłowska, Krzysztof Kowalski, Halina Woźniak, Alicja Koszela, Maciej Adamowicz, Marek Cierpień-Wolan, Remigiusz Tworzydło	
Wspomnienie o Władysławie Wiesławie Łagodzińskim	62
Remembering Władysław Wiesław Łagodziński	
Dyskusje. Recenzje. Informacje	
Discussions. Reviews. Information	
Justyna Gustyn	
Wydawnictwa GUS. Luty 2021	80
Publications of Statistics Poland. February 2021	
Dla autorów	83
For the authors	
Zakres tematyczny działów	94
Thematic scope of sections	

OD REDAKCJI

W marcowym wydaniu „Wiadomości Statystycznych. The Polish Statistician” proponujemy Państwu lekturę artykułów z zakresu studiów metodologicznych, praktycznych zastosowań statystyki i studiów interdyscyplinarnych.

Prof. dr hab. Mirosław Szreder w artykule *Błędy nielosowe i ich znaczenie w testowaniu hipotez* omawia znaczenie błędów nielosowych w podejmowaniu decyzji opartych na wykorzystaniu klasycznej procedury weryfikacji hipotez. Uzasadnia twierdzenie, że w przypadku dużych prób testy statystyczne stają się wrażliwsze na oddziaływanie błędów nielosowych, i dowodzi, że błędy systematyczne zwiększają prawdopodobieństwo odrzucenia prawdziwej hipotezy, gdy liczebność próby wzrasta. Wskazuje, że jednym ze sposobów zmniejszenia ryzyka błędnych decyzji w testowaniu hipotez może być analiza przedziału ufności przeprowadzona oprócz weryfikacji hipotez lub zamiast niej.

W badaniu przedstawionym w artykule *Czynniki determinujące dysproporcje w wynagrodzeniach kobiet i mężczyzn w krajach Unii Europejskiej* prof. dr hab. Dorota Witkowska oraz dr Aleksandra Matuszewska-Janica weryfikują rozmiar nieskorygowanej luki płacowej pracowników najemnych w Unii Europejskiej po kryzysie zapoczątkowanym w 2007 r. w odniesieniu do sytuacji sprzed kryzysu. Autorki wskazują też zmiany w poziomie zatrudnienia kobiet i mężczyzn oraz w zróżnicowaniu wynagrodzeń ze względu na płeć, mierzonym wskaźnikiem *gender pay gap*. W analizie wykorzystują wyniki badań Structure of Earnings Survey oraz Labour Force Survey, udostępniane przez Eurostat. Stwierdzają, że w 24 krajach UE nastąpił wzrost feminizacji zatrudnienia, a luka płacowa w latach 2006–2018 nie uległa istotnemu zmniejszeniu.

Wykorzystanie Google Trends do modelowania stopy bezrobocia rejestrowanego w Polsce to temat artykułu Mariusza Malinowskiego. Metoda zastosowana w badaniu opiera się na technikach nowcastingu służących do oceny bieżącego stanu gospodarki. Autor przedstawia zarówno możliwości, jak i ograniczenia wykorzystywania nowego źródła danych, jakim jest serwis Google Trends (w literaturze anglojęzycznej już dość powszechnego), w analizach makroekonomicznych dotyczących Polski. Na podstawie danych za lata 2004–2019 pochodzących z badań GUS oraz Google Trends stwierdza, że indeksy Google nie poprawiają trafności predykcji modelu autoregresyjnego. Zadowalające rezultaty uzyskuje się jedynie w przypadku indeksów związanych z międzynarodową mobilnością siły roboczej.

W dziale In memoriam wspominamy Władysława Wiesława Łagodzińskiego. Współpracownicy i przyjaciele ukazują jego drogę zawodową – ponad 40 lat pracy w GUS na odpowiedzialnych stanowiskach, pracę w Urzędzie Statystycznym w Warszawie, aktywną działalność w Polskim Towarzystwie Statystycznym, inicjowanie i współrealizowanie badań statystycznych, uczestnictwo w Kolegium Redakcyjnym „WS”. Dzieli się także osobistymi wspomnieniami związanymi z jego osobą. Sylwetkę i dokonania Władysława Wiesława Łagodzińskiego dopełnia wykaz najważniejszych publikacji.

Wydanie kończy się omówieniem nowości wydawniczych GUS przygotowanym przez Justynę Gustyn.

Zapraszamy do lektury.

FROM THE EDITORIAL TEAM

In the March issue of *Wiadomości Statystyczne. The Polish Statistician*, the Reader can find articles from the realm of methodological studies, practical applications of statistics and interdisciplinary studies.

Mirosław Szreder, PhD, DSc, ProfTIt, in his article entitled *Non-random errors and their importance in testing of hypotheses* takes a closer look at the role non-random errors play in making decisions based on the classical method of hypothesis verification. The author proves that statistical tests become more sensitive to the effects of non-random errors when they are performed on large samples; moreover, systematic errors tend to increase the probability of the rejection of a true hypothesis as the size of the sample grows. He also indicates that what could be one of the ways to limit the risk of making wrong decisions in the course of testing hypotheses, is the analysis of confidence intervals carried out either parallel to or instead of the verification of a hypothesis.

In the study presented in the article *Factors determining disproportions in men and women's wages in the European Union countries* by Dorota Witkowska, PhD, DSc, ProfTIt, and Aleksandra Matuszewska-Janica, PhD, the authors verify the volume of the unadjusted gap in wages of hired workers in the EU after the 2007 crisis compared to the pre-crisis situation. The authors also point to changes in the employment rates of men and women, as well as to changes in the pay gap between the two sexes, measured by the gender pay gap indicator. In their analysis, the authors utilize the results of Eurostat's Structure of Earnings Survey and the Labour Force Survey. According to the study, the employment in 24 EU countries has become increasingly 'feminised' since the end of the crisis, whereas the pay gap did not significantly shrink in the period of 2006–2018.

The study described in Mariusz Malinowski's *Use of Google Trends in modelling registered unemployment rate in Poland* is carried out according to a method utilising nowcasting techniques, commonly used for the assessment of the state of the economy. The author presents both the opportunities connected to the adoption of a new source of data – Google Trends (already relatively popular in the literature written in English), and its limitations, as they arise in the process of performing macroeconomic analyses concerning Poland. On the basis of data for 2004–2019 taken from Statistics Poland's research and those retrieved from Google Trends, the author draws a conclusion that Google indexes do not improve the accuracy of prediction of autoregressive models. Satisfactory results are only obtained in the case of indexes measuring the international mobility of workers.

The In memoriam section features the life and achievements of Władysław Wiesław Łagodziński. His colleagues and friends present the course of his professional career – over 40 years of work for Statistics Poland on several important positions, work for the Statistical Office in Warsaw, active membership in the Polish Statistical Association, initiating and co-supervising statistical research, and inspiring participation in the work of the WS Editorial Board. They also share with the Reader their personal memories of the statistician. The tribute to Władysław Wiesław Łagodziński concludes with the list of his most important publications.

The issue closes with Justyna Gustyn's presentation of Statistics Poland's new publications. We wish you pleasant reading.

Błędy nielosowe i ich znaczenie w testowaniu hipotez

Mirosław Szreder^a

Streszczenie. We współczesnych badaniach reprezentacyjnych coraz częściej dają o sobie znać błędy o charakterze nielosowym, w tym w szczególności wynikające z braków odpowiedzi lub źle wykonanych pomiarów (niedokładnej obserwacji statystycznej). Do tej pory rzadko dysutowano o skutkach tego typu błędów w procedurze weryfikacji hipotez statystycznych. Uwaga badaczy skupiała się niemal wyłącznie na błędzie losowania (błędzie losowym). Błąd ten maleje wraz ze wzrostem liczebności próby. To sprawia, że badacze, nierzadko mający do dyspozycji bardzo duże liczebnie próby, tracą z pola widzenia konsekwencje nie tylko błędu losowego, lecz także błędów nielosowych.

Celem artykułu jest wskazanie na znaczenie błędów nielosowych w podejmowaniu decyzji opartych na wykorzystaniu klasycznej procedury weryfikacji hipotez. Szczególną uwagę poświęcono sytuacjom, w których badacz dysponuje dużą liczebnie próbą. W pracy uzasadniono twierdzenie, że w dużych próbach testy statystyczne stają się bardziej wrażliwe na oddziaływanie błędów nielosowych. Błędy systematyczne, będące szczególnym przypadkiem błędów nielosowych, zwiększają prawdopodobieństwo błędnej decyzji o odrzuceniu prawdziwej hipotezy wraz ze wzrostem liczebności próby. Wzbogacenie weryfikacji hipotez o analizę opartą na estymacji przedziałowej może wspomóc badacza w poprawnym wnioskowaniu.

Słowa kluczowe: testowanie hipotez statystycznych, błąd losowania, błąd losowy, błędy nielosowe

JEL: C12, C13, C18, D80

Non-random errors and their importance in testing of hypotheses

Abstract. Increasing numbers of non-random errors are observed in contemporary sample surveying – in particular, those resulting from no response or faulty measurements (imprecise statistical observation). Until recently, the consequences of these kinds of errors have not been widely discussed in the context of the testing of hypotheses. Researchers focused almost entirely on sampling errors (random errors), whose magnitude decreases as the size of the random sample grows. In consequence, researchers who often use samples of very large sizes tend to overlook the influence random and non-random errors have on the results of their study.

The aim of this paper is to present how non-random errors can affect the decision-making process based on the classical hypothesis testing procedure. Particular attention is devoted to cases in which researchers manage samples of large sizes. The study proved the thesis that samples of large sizes cause statistical tests to be more sensitive to non-random errors. Systematic errors, as a special case of non-random errors, increase the probability of making the wrong decision to reject a true hypothesis as the sample size grows. Supplementing the testing of hypotheses with the analysis of confidence intervals may in this context provide substantive support for the researcher in drawing accurate inferences.

Keywords: testing of hypotheses, sampling error, random error, non-random errors

^a Uniwersytet Gdański, Wydział Zarządzania, Katedra Statystyki / University of Gdansk, Faculty of Management, Department of Statistics. E-mail: mirosław.szreder@ug.edu.pl.
ORCID: <https://orcid.org/0000-0002-7597-0816>.

1. Wprowadzenie

Rzadko które metody i techniki statystyczne budzą w środowisku naukowym tyle kontrowersji, co statystyczne testy istotności (Dorofeev i Grant, 2006). Źródła nieporozumień i polemik jest tu co najmniej kilka. Wiążą się one zarówno z niejednokowym rozumieniem w środowisku statystyków i poza nim kluczowego w testowaniu hipotez terminu *istotność*, jak i z brakiem konsensusu wśród części naukowców co do tego, w jakich warunkach mogą być stosowane testy i jak poprawnie należy interpretować ich rozstrzygnięcia (szerzej: Szreder, 2019).

Retrospektywne spojrzenie na ewolucję zastosowań testów statystycznych pozwala dostrzec, że wykorzystywane były one do realizacji dwóch odrębnych celów (Gigerenzer i Marewski, 2015). Sensem pierwszego – wcześniejszego z zastosowań, obecnego m.in. w fizyce i astronomii – było zaufanie do pewnej hipotezy, np. że rozkład błędów pomiarów wobec prawdziwego położenia określonej planety jest rozkładem normalnym, i odrzucenie obserwacji, które nazbyt odbiegały od tego rozkładu (najczęściej z winy obserwatora lub urządzeń pomiarowych). Celem tego rodzaju testowania było odrzucenie nie hipotezy, lecz niektórych obserwacji odstających (*outliers*). Drugie z zastosowań jest w swojej istocie niejako odwróceniem pierwszego. Uzyskane w próbie obserwacje uznaje się za godne zaufania, a celem testowania jest odrzucenie hipotezy, która okaże się zbyt odległa od tego, co w rzeczywistości zaobserwowano.

W klasycznym paradygmacie testowania hipotez statystycznych, zarówno w ujęciu Fishera, jak i Neymana-Pearsona, obserwacje uzyskane w próbie losowej nie podlegają ocenie ani weryfikacji, lecz służą stwierdzeniu, na ile nieprzystająca do tych obserwacji jest testowana hipoteza. Stopień nieprzystawania tej hipotezy (zerowej) do zaobserwowanej próby jest wyrażany w kategoriach probabilistycznych i stanowi o tym, czy hipoteza ta zostanie uznana za nieprawdziwą, a w konsekwencji – odrzucona. Najważniejszą rolę w rozstrzyganiu o nieprawdziwości hipotezy odgrywa więc próba statystyczna. A jeżeli tak, to w całej procedurze znaczenie mają zarówno błąd losowania, jak i błędy w dokonanych pomiarach, stanowiące jedną z kategorii błędów o charakterze nielosowym. O występowaniu i znaczeniu błędów nielosowych zdaje się zapominać część współczesnych badaczy i użytkowników metod statystycznych, którzy kładą nacisk niemal wyłącznie na pożądane wartości wskaźnika *p-value* i dążenie do uzyskania statystycznie istotnego rozstrzygnięcia. Trzeba podkreślić, że postawa twórców wnioskowania statystycznego była w tym zakresie zgoła odmienna. William S. Gosset – twórca rozkładu *t*-Studenta – uważał małe błędy pomiaru za ważniejsze od małych wartości *p-value* (Gigerenzer i in., 1989).

W naukach ekonomicznych i społecznych kategoria błędu pomiaru jest szczególnie pojemna, ponieważ obejmuje szereg uświadomionych i nieuświadomionych zniekształceń i pomyłek leżących zarówno po stronie realizatorów badania, w tym ankieterów, jak i po stronie respondentów. Nielosowy charakter, skutkujący często systematycznym obciążeniem, mają także inne błędy obecne w badaniach reprezentacyjnych, w tym błędy pokrycia i błędy braków odpowiedzi¹. Warto więc rolę tych błędów oraz innych ułomności składających się na pojęcie błędów nielosowych rozważyć w kontekście współczesnego kryzysu replikowalności eksperymentów (w naukach ścisłych i społecznych), a szerzej – kryzysu statystycznej istotności².

Celem artykułu jest wskazanie na znaczenie błędów nielosowych w podejmowaniu decyzji opartych na wykorzystaniu klasycznej procedury weryfikacji hipotez. Szczególną uwagę poświęcono sytuacjom, w których badacz dysponuje dużą liczebnie próbą badawczą. Duża liczba obserwacji w próbie bywa współcześnie uznawana za usprawiedliwienie niespełnienia w praktyce empirycznej niektórych warunków wnioskowania statystycznego albo za remedium na ten stan rzeczy. Warto więc pokazać, jakie są skutki błędów losowych i nielosowych w testowaniu hipotez dla dużych prób. Do tej pory rzadko rozważano te konsekwencje w odniesieniu do błędów nielosowych.

Rozpocząć jednak trzeba od przypomnienia warunków, jakie powinny być spełnione, aby można było poprawnie stosować metody wnioskowania statystycznego, w tym weryfikacji hipotez. Najbardziej bowiem kontrowersyjną kwestią jest dziś coraz śmielsze dążenie do posługiwania się technikami wnioskowania statystycznego i interpretowanie jego wyników w kategoriach probabilistycznych, kiedy obserwacje w próbie nie mają charakteru losowego.

2. Znaczenie założenia o losowości próby

Łatwiejszy niż w przeszłości dostęp do danych statystycznych oraz do specjalistycznego oprogramowania komputerowego sprawia, że w wielu dziedzinach poznania naukowego wzrosło w ostatnich latach zainteresowanie metodami opisu statystycznego oraz wnioskowania. W szczególności testy statystyczne stały się jedną z najważniejszych grup narzędzi badawczych w różnych obszarach zastosowań nauk empirycznych. Wykorzystywane są do rozstrzygania na podstawie obserwacji w próbie o prawdziwości lub nieprawdziwości sformułowanych wcześniej hipotez odnoszących się do charakterystyk badanej populacji. Powszechność zastosowań technik wnioskowania statystycznego, w tym weryfikacji hipotez, świadczy z jednej strony

¹ O występowaniu tych kategorii błędów w badaniach statystyki publicznej pisze m.in. Paradysz (1989, 2009).

² Temat tego kryzysu obszernie – w ponad 40 artykułach – opisali statystycy z całego świata na łamach czasopisma „The American Statistician” w nr. 73 z marca 2019 r.

o potencjale i możliwościach tych technik, a z drugiej – przy masowym ich wykorzystaniu – staje się powodem uzasadnionych obaw o poprawność ich użycia. W literaturze światowej najszerzej w tym kontekście dyskutowana jest kwestia rozstrzygnięcia za pomocą wskaźnika *p-value* o nieprawdziwości testowanej hipotezy, a także stosowania w tej procedurze stwierdzenia o *statystycznej istotności* efektu zaobserwowanego w próbie³. Jednak pierwotne, czyli wyjściowe, wobec tych kontrowersji jest inne zagadnienie, mianowicie sens wymogu losowości próby – jednego z najbardziej podstawowych założeń w teorii wnioskowania statystycznego.

Matematycznym modelem wnioskowania statystycznego jest – przypomnijmy – model probabilistyczny, który opiera się na założeniu o losowym mechanizmie generowania obserwacji w próbie. Z tym założeniem związane są najbardziej podstawowe i najważniejsze elementy estymacji i weryfikacji hipotez. Punktem wyjścia do zastosowań probabilistycznego modelu wnioskowania statystycznego jest uznanie, że nawet próba losowa nie stanowi doskonałego odzwierciedlenia struktury populacji, a w powtarzalnych losowaniach mogą występować znaczne różnice w charakterystykach otrzymywanych kolejno prób. Innymi słowy, jest to cechująca badacza świadomość istnienia błędu losowania (*sampling error*), nazywanego także błędem losowym (*random error*), który wyraża niemożność zapewnienia przez mechanizm generujący losowe obserwacje w próbie pełnej zgodności struktury próby ze strukturą populacji. Duże znaczenie tego błędu – który nierozzerwalnie wiąże się z samym aktem losowania – we wnioskowaniu statystycznym wynika przynajmniej z trzech powodów.

Po pierwsze błąd losowania stanowi jedyny rodzaj błędu, który występuje w teoretycznym (matematycznym) modelu wnioskowania statystycznego. Brak w tym modelu innych elementów całkowitego błędu badania próbkowego (*total survey error*). W szczególności nieobecne są w nim błędy o charakterze nielosowym.

Po drugie błąd losowania – wraz z założeniami o rozkładach prawdopodobieństwa cech populacji, z której pochodzą losowe obserwacje – stanowi punkt wyjścia do określenia wszelkich miar precyzji i wiarygodności wnioskowania statystycznego. Dodatkowo od tego właśnie błędu (i tylko tego) rozpoczyna się procedura wyprowadzania formuł matematycznych i konstruowania narzędzi wnioskowania, które odgrywają kluczową rolę w uogólnianiu prawidłowości zaobserwowanych w próbie. Należą do nich przedziały ufności czy formuły na minimalną liczebność próby oraz najważniejsze statystyki związane z testowaniem hipotez, w tym *p-value*.

³ Oby tym kwestiom poświęcają coraz więcej uwagi prestiżowe czasopisma naukowe, m.in.: „Science” (zob. Amrhein, Greenland i McShane, 2019) oraz „The American Statistician” (zob. Wasserstein i Lazar, 2016). Niektórzy autorzy piszą wręcz o kulcie istotności statystycznej, przy czym zaznaczają już na wstępie, że istotność statystyczna nie musi implikować istotności w podstawowym tego słowa znaczeniu, a niekiedy wręcz nie oznacza niczego, co byłoby warte uwagi (oryg. „Statistically significant relationships may, and often do, tell us nothing that matters”), o czym pisze Lempert (2009, s. 226). Zob. także: Gelman i Stern (2006), Ziliak i McCloskey (2008) oraz Szreder (2019).

Po trzecie oceny i decyzje mają we wnioskowaniu statystycznym właściwą, poprawną logicznie interpretację tylko wtedy, gdy odnoszą się do hipotetycznego ciągu powtarzalnych prób losowych i generowanego błędu losowania. Dotyczy to zarówno interpretacji przedziałów ufności, błędów pierwszego i drugiego rodzaju w testowaniu hipotez oraz wskaźnika *p-value*, jak i wszelkich właściwości estymatorów. Jeżeli badacz nie potrafi uzasadnić, że próba, którą dysponuje, jest losowa albo może być traktowana jak losowa, to wspomniane narzędzia wnioskowania statystycznego przestają zapewniać poprawną interpretację⁴ (Hirschauer i in., 2020, s. 72)⁵. Podobnie dzieje się, gdy obserwacją są objęte wszystkie jednostki populacji, co w całości eliminuje błąd losowania w analizie statystycznej⁶.

Wykorzystanie rachunku prawdopodobieństwa we wnioskowaniu statystycznym stwarza możliwość matematycznego opisu relacji pomiędzy próbą losową a populacją, którą ta próba reprezentuje. W naukach społecznych i eksperymentalnych za próbę losową stanowiącą podstawę wnioskowania uznaje się albo losowo uzyskany podzbiór jednostek ze zdefiniowanej wcześniej populacji (do której odnoszone będą późniejsze wnioski), albo zrandomizowany kontrolowany eksperyment, w którym jednostki poddane obserwacji trafiają losowo do dwóch grup: eksperymentalnej (poddanej oddziaływaniu jakiegoś bodźca) i kontrolnej⁷.

Wydaje się, że jedną z głównych przyczyn niewystarczającej koncentracji badaczy na źródłach i konsekwencjach błędów nielosowych, w tym w szczególności błędów pomiaru, jest mylne przekonanie, że w warunkach łatwej dostępności dużych prób problem ten zanika albo wręcz nie istnieje. Duża liczebność próby zwodzi niektórych badaczy, upatrujących w niej gwaranta poprawnego i wiarygodnego wnioskowania niezależnie od jakości obserwacji w próbie. Najczęściej zapominają oni o tym, że najważniejsze kategorie błędów nielosowych nie są funkcjami liczebności próby. I nie tylko dlatego warto przeanalizować wzajemne relacje między liczebnością próby i błędami nielosowymi z jednej a poprawnością testowania hipotez z drugiej strony. Gigerenzer i Marewski (2015, s. 425) stwierdzają, że proste zwiększanie liczby

⁴ Oryg. „they become essentially uninterpretable”.

⁵ Zob. także Vogt i in. (2014, s. 244). Autorzy stwierdzają, że w badaniach, w których nie korzysta się z randomizacji lub z próbkowania losowego, klasyczne podejście do wnioskowania statystycznego jest nieodpowiednie (oryg. „in research not employing random assignment or random sampling, the classical approach to inferential statistics is inappropriate”).

⁶ W tego typu sytuacjach (tj. w badaniach wyczerpujących) swoje znaczenie traci także termin *istotność statystyczna* w odniesieniu do parametrów populacji, np. różnic tych parametrów od zera. Źródłem tego terminu jest bowiem świadomość badacza, że w każdej zaobserwowanej próbie losowej tkwią konsekwencje błędu losowania. Statystycznie istotne różnice to takie, które nie dają się wyjaśnić oddziaływaniem wyłącznie błędu losowania, a ściślej – czyniłoby takie wyjaśnienie mało wiarygodnym (prawdopodobnym) w świetle dokonanych obserwacji. To uściślenie, dotyczące użycia prawdopodobieństwa w ocenie istotności statystycznej, jest szczególnie potrzebne w kontekście uzasadnionej, szerokiej krytyki dychotomizacji pojęcia *istotności statystycznej* – traktowania tego pojęcia zero-jedynkowo (binarnie).

⁷ Gdy ściśła randomizacja jest niemożliwa w danym eksperymencie, zaleca się użycie takich technik, jak *propensity score* w celu zmniejszenia błędu wyboru (*selection bias*). Zob. np. Mercer i in. (2017).

obserwacji w próbie stało się środkiem zastępczym (substytutem) dążenia do minimalizowania błędów⁸. O ile jednak błąd losowania maleje wraz ze wzrostem liczby obserwacji w próbie, o tyle ważne jest, aby prawidłowości tej nie odnosić do całkowitego błędu wnioskowania, ponieważ jego pozostałe składniki, mające charakter nie-losowy, nie podlegają tej prawidłowości.

Niedocenianiu błędów nielosowych sprzyja nie tylko złudne przekonanie o możliwościach sprawczych dużych prób, lecz także bardziej złożone zjawisko pomijania i lekceważenia wielu innych okoliczności w dążeniu do szybkiego osiągnięcia statystycznie istotnego wyniku. Ma to związek z błędnym wyobrażeniem odbiorców badań, a także niektórych naukowców, że wartościowe są tylko te wyniki, którym można przypisać etykietę statystycznie istotnych. W dążeniu do uzyskania statystycznie istotnych wyników pomija się zaś kwestię błędów nielosowych. Tymczasem te ostatnie mogą silniej od błędu losowania wpływać na rozstrzygnięcia o istotności różnic stanowiących podstawę do odrzucenia hipotezy zerowej.

3. Testowanie hipotez na podstawie liczebnie dużych prób

Problem błędów nielosowych w testowaniu hipotez nie ujawniłby się prawdopodobnie z taką ostrością, gdyby nie to, że coraz częściej badacze mają dostęp do liczebnie dużych prób, znacznie przekraczających wyobrażenia twórców teorii wnioskowania statystycznego z początków XX w. Wykorzystanie wielotysięcznej próby do weryfikacji hipotez statystycznych rodzi różne konsekwencje, z których tylko część była do tej pory wyraźnie komunikowana. Najczęściej skupiano się na poprawie mocy testu wraz ze wzrostem liczebności próby. Innymi słowy, stwierdzano słusznie, że rosnąca liczebność próby zwiększa zdolność testu statystycznego do prawidłowego rozróżnienia pomiędzy hipotezą prawdziwą a fałszywą. Raz jeszcze warto zaznaczyć, że moc testu (dopełnienie do 1 prawdopodobieństwa błędu drugiego rodzaju) nie uwzględnia żadnych innych – poza błędem losowania – błędów, którymi mogą być obciążone wyniki w próbie.

Pozostańmy na razie przy powyższym założeniu, abstrahując od oddziaływania błędów nielosowych. Już tutaj można dostrzec pierwsze wyzwania, przed jakimi staje statystyk dysponujący dużą próbą losową. Warto zaznaczyć, że po raz pierwszy wskazywano na nie już ponad pół wieku temu. W 1966 r. na łamach „Psychological Bulletin” David Bakan trafnie uzasadniał, że każda hipoteza zerowa może zostać odrzucona, jeżeli tylko wylosuje się odpowiednio dużą próbę. Zilustrował to m.in. doświadczeniem, w którym sam uczestniczył, polegającym na przetestowaniu zwykłych baterii przez 60 tys. Amerykanów. Za pomocą podziału tej próby na dwie gru-

⁸ Oryg. „Simply increasing the number N of subjects became a surrogate for minimizing errors”.

py osób dokonywanego według zupełnie dowolnych i nic nieznaczących kryteriów, takich jak mieszkanie po wschodniej lub zachodniej stronie rzeki Missisipi, na północy lub południu kraju, uzyskiwał za każdym razem statystycznie istotne różnice w średnich. I nie były to w tamtym czasie jedyne spostrzeżenia prowadzące do wniosku, że „jeżeli hipoteza zerowa nie zostaje odrzucona, to zwykle z tego powodu, że liczebność próby (n) była za mała”⁹ (Nunnally, 1960, s. 643). W znanej monografii z końca lat 70. XX w. Edward E. Leamer pisał bardziej dosadnie, że w dużych próbach nawet bardzo mała i w danym zagadnieniu nieznacząca wielkość efektu będzie prowadzić do stwierdzenia jego statystycznej istotności (czyli odrzucenia hipotezy zerowej), „ponieważ duża próba jest przypuszczalnie bardziej miarodajna niż mała, a hipotezę zerową odrzucimy dla odpowiednio dużej próby, moglibyśmy równie dobrze tę hipotezę odrzucić bez pobierania próby w ogóle”¹⁰ (1978, s. 89).

Opisane wyżej prawidłowości można uzasadnić także analitycznie, odwołując się do najbardziej popularnego obecnie wskaźnika w procedurze weryfikacji hipotez – *p-value*. Lin i in. (2013) prezentują matematyczny dowód na to, że wartości *p-value*, oparte na zgodnych estymatorach, charakteryzują się następującą własnością asymptotyczną w odniesieniu do hipotezy zerowej $H_0: \beta = 0$:

$$\lim_{n \rightarrow \infty} p\text{-value}(n) = \lim_{n \rightarrow \infty} P_n(|\hat{\beta} - \beta| < \varepsilon) = \begin{cases} 0 & \text{dla } \beta \neq 0 \\ 1 & \text{dla } \beta = 0 \end{cases}$$

Oznacza to, że przy wzrastającej do nieskończoności liczebności próby (n) cała masa prawdopodobieństwa w rozkładzie próbkowym zgodnego estymatora $\hat{\beta}$ coraz ściślej skupia się wokół parametru populacji β . Jeżeli w populacji parametr β jest dokładnie równy hipotetycznej wartości 0, wówczas dla dużych liczebności prób rozkład estymatora $\hat{\beta}$ jest tak silnie skoncentrowany wokół 0, że bliskie pewności jest otrzymanie próby, w której $\hat{\beta} = 0$. To z kolei oznacza – przy wzięciu pod uwagę interpretacji *p-value* – że w tej sytuacji wraz ze wzrostem liczebności próby wartości *p-value* rosną do 1. W pozostałych przypadkach, kiedy parametr populacji odbiega od 0 (od wartości hipotetycznej), choćby bardzo nieznacznie (np. na dalekim miejscu po przecinku), to konsekwentnie trzeba zauważyć, że właśnie w tym punkcie – bliskim, ale nie równym 0 – skupia się niemal całe prawdopodobieństwo rozkładu próbkowego estymatora $\hat{\beta}$. To oznacza, że dla $n \rightarrow \infty$ zbiega do 1 prawdopodobieństwo otrzymania próby, dla której $\hat{\beta} = \beta \neq 0$, a do 0 – prawdopodobieństwo, które wyraża wskaźnik *p-value* – otrzymania zaobserwowanej próby przy założeniu, że

⁹ Oryg. „If the null hypothesis is not rejected, it usually is because the N is too small”. Tłumaczenie tego i wszystkich pozostałych cytatów w artykule pochodzi od autora.

¹⁰ Oryg. „A large sample is presumably more informative than a small one, and since it is apparently the case that we will reject the null hypothesis in a sufficiently large sample, we might as well begin by rejecting the hypothesis and not sample at all”.

prawdziwa jest hipoteza zerowa $\beta = 0$. Innymi słowy, z wyjątkiem rzadkich w praktyce badawczej sytuacji, kiedy parametr β jest dokładnie – co do dalekich miejsc po przecinku – równy wartości ujętej w hipotezie zerowej (wtedy p -value zbiega do 1), wartości wskaźnika p -value są zbieżne do 0 przy założeniu rosnącej do nieskończoności liczebności próby.

W realnych populacjach poddawanych analizom i wnioskowaniu testowany parametr bardzo rzadko bywa dokładnie równy tej wartości, jaką zapisano w hipotezie zerowej. Dlatego liczebnie duże próby będą prawie zawsze dawały wartości p -value bliskie 0, skutkujące odrzuceniem hipotezy zerowej. Z wyraźną łatwością odrzucane są np. hipotezy o tym, że współczynniki korelacji czy regresji są równe 0, jeżeli tylko badacz posiada odpowiednio dużą próbę losową. W populacji bowiem niesłuchanie rzadko wartości tych współczynników są dokładnie równe 0 (jak głosi hipoteza zerowa). Badacz mający do dyspozycji dużą próbę bez trudu odrzuci hipotezę głoszącą, że testowany parametr jest statystycznie nieistotny (równy 0), mimo że w rzeczywistości jego wartość będzie mogła być bardzo bliska 0. Prowadzi to oczywiście do przykrych błędów poznawczych – rozbieżności między istotnością naukową a istotnością statystyczną. Tę kłopotliwą konkluzję dla zastosowań testów w dużych próbach wyraził dobitnie m.in. Cohen (1990, s. 1308):

Hipoteza zerowa traktowana dosłownie (a tylko tak można ją rozumieć w formalnej procedurze weryfikacji hipotez) jest zawsze fałszywa w realnej rzeczywistości... Jeżeli zaś jest fałszywa, choćby w bardzo nieznacznym stopniu, to z odpowiednio dużej próby uzyska się wynik istotny statystycznie, który doprowadzi do jej odrzucenia¹¹.

Ta decyzja z kolei oznacza, że dla dużych prób wiele nieznaczących różnic i określonych w hipotezie zerowej relacji między zmiennymi zostanie uznanych za istotne statystycznie. Według Lemperta (2009, s. 230) „przy odpowiednio dużej wielkości próby niemal wszystkie zależności w próbie będą istotne statystycznie, ponieważ dla rosnącej liczebności próby znikają efekty losowe i nawet słabe sygnały przebijają się przez szum realnego świata”¹². Przez *sygnał* należy rozumieć wartość statystyki testowej w próbie, wyrażającą niedopasowanie uzyskanej próby do brzmienia hipotezy zerowej. Natomiast desygnatem *szumu* – albo zakłóceń – jest wartość miary rozproszenia statystyki testowej, np. odchylenie standardowe. W dużych próbach, jak wskazano wcześniej, należy oczekiwać niewielkich zakłóceń (małych wartości odchy-

¹¹ Oryg. „The null hypothesis, taken literally (and that’s the only way you can take it in formal hypothesis testing), is always false in the real world... If it is false, even to a tiny degree, it must be the case that a large enough sample will produce a significant result and lead to its rejection”.

¹² Oryg. „With a large enough N, virtually all associations in a sample will be statistically significant, for as size increases, random effects are more likely to cancel out, and even weak signals will emerge through the real world’s noise”.

lenia standardowego statystyki testowej)), stąd nawet niewielka niezgodność próby z wynikami, jakie powinny się były pojawić, gdyby prawdziwa była hipoteza zerowa, doprowadzi do odrzucenia tej hipotezy.

W tym kontekście rzadko się do tej pory zwracało uwagę na błędy o charakterze nielosowym, których oddziaływanie dodatkowo zwiększa szanse odrzucenia hipotezy zerowej w dużych próbach. Dotyczy to w szczególności sytuacji, gdy testowana hipoteza (np. o statystycznie nieistotnej różnicy parametru populacji od 0) nie powinna zostać uznana za nieprawdziwą i odrzucona. Znaczenie tej kategorii błędów (nielosowych) wyraźnie wzrosło we współczesnych badaniach statystycznych¹³, co wciąż pozostaje niedocenione w światowej literaturze statystycznej i w praktyce badań empirycznych. Efektem oddziaływania błędów nielosowych jest najczęściej obciążenie uzyskanych ocen błędem systematycznym. Co prawda wielkość tego błędu nie jest funkcją wielkości próby, ale jego wpływ na poprawność wnioskowania (słuszność odrzucenia testowanej hipotezy) wykazuje taką zależność.

Gdy systematyczne obciążenie jest nieduże, a liczebność próby niewielka, to ryzyko błędnej decyzji o odrzuceniu hipotezy zerowej, gdy jest ona prawdziwa, nie jest duże. Przed tego typu błędnymi decyzjami chroni duże rozproszenie statystyki testowej, wyrażające znaczną niepewność badacza. Obszar nieodrzuconia hipotezy zerowej jest w tym przypadku spory. Ryzyko to oznacza, w kontekście wskaźnika *p-value*, że gruby ogon w rozkładzie statystyki testowej czyni dość prawdopodobnym przyjęcie przez nią wartości odległych od tej, która wynikałaby z hipotezy zerowej (co odpowiada stosunkowo dużym wartościom *p-value*). W rezultacie hipoteza zerowa nie zostanie odrzucona.

Systematyczny błąd staje się groźniejszy wtedy, gdy towarzyszy mu duża liczebność próby. Statystyczna prawidłowość, silnie potwierdzona przez dużą wielkość próby – a w rezultacie przez małe rozproszenie statystyki testowej – jest trudna do zakwestionowania. Rzadko się jednak podkreśla, że każda taka prawidłowość jest w dużym stopniu warunkowa, ponieważ zachodzi jedynie wtedy, gdy obserwacje w próbie nie są obciążone błędami nielosowymi, w szczególności systematycznymi. Istnienie błędu systematycznego w dokonanych pomiarach w próbie – przy małym lub bardzo małym rozproszeniu statystyki testowej – sprawia, że obszar odrzucenia hipotezy zerowej znacznie się rozszerza. Innymi słowy, jakakolwiek wartość statystyki różna od tej, która wynika z brzmienia hipotezy zerowej, staje się bardzo mało prawdopodobna. Uzyskanie takiej lub bardziej skrajnej wartości (*p-value*) jest bliskie 0. Błąd systematyczny będzie w tego typu sytuacjach zwiększał ryzyko niesłusznego odrzucenia hipotezy zerowej. Im większa jest liczebność próby, tym bardziej wrażliwy na błędy systematyczne staje się test statystyczny. Co prawda niewielki błąd tego

¹³ O zmianach w strukturze całkowitego błędu badania próbkowego pisze m.in. Szreder (2015).

rodzaju oznacza niewielkie przesunięcie rozkładu statystyki testowej względem osi rzędnych, ale przy bardzo małej dyspersji rozkładu tej statystyki ogromna część masy prawdopodobieństwa przesuwa się w kierunku nieprawdziwej wartości. W rozkładach bardziej rozproszonych, odpowiadających mniejszym liczebnościom próby, jest to mniejsza frakcja masy prawdopodobieństwa pod krzywą rozkładu. Zależności te ilustruje poniższy przykład.

Załóżmy, że weryfikacji na poziomie istotności 0,05 podlega hipoteza stwierdzająca, że średnia μ w populacji o rozkładzie normalnym, w którym odchylenie standardowe wynosi 4, równa jest 5:

$$H_0: \mu = 5$$

$$H_1: \mu \neq 5$$

Oznacza to, że jeżeli rzeczywiście średnia wartość w tej populacji wynosi 5, to hipotezę zerową odrzucimy w tych próbach losowych, dla których wartość statystyki testowej

$$Z = \frac{\bar{X} - \mu}{\sigma} \cdot \sqrt{n}$$

przyjmie wartości mniejsze od $-1,96$ lub większe od $+1,96$, gdzie $\bar{X}$ oznacza średnią arytmetyczną z próby, σ – odchylenie standardowe w populacji (równe 4), a n – liczebność próby. W tych przypadkach zmaterializuje się błąd pierwszego rodzaju (odrzućenie prawdziwej hipotezy), a prawdopodobieństwo jego popełnienia wyniesie 0,05. Innymi słowy, jeżeli wnioskowanie wolne jest od innych błędów – w szczególności nielosowych – to przeciętnie 5 na 100 wylosowanych prób będzie prowadzić do błędnej decyzji o odrzuceniu prawdziwej hipotezy zerowej.

Z kolei aby zilustrować skutki błędów nielosowych – a zwłaszcza konsekwencje błędu systematycznego – można obliczyć, jak rośnie prawdopodobieństwo podjęcia błędnej decyzji o odrzuceniu prawdziwej hipotezy zerowej, gdy obserwacje w próbie obciążone są takim właśnie błędem. Prawdopodobieństwo to jest wyraźnie większe od 0,05 i rośnie wraz ze wzrostem wielkości próby. Obliczenia obrazujące, jak się ono zmienia, wykonano dla dwóch wielkości błędu systematycznego, wynoszącego:

- a) $+0,5$ (wszystkie dokonane pomiary w próbie są zawyżone o 0,5 jednostki pomiaru);
- b) $+1,0$ (wszystkie pomiary są zawyżone o 1 jednostkę pomiaru).

Dla ustalonych liczebności prób $n = 16, 36, 80, 100, 200, 500$ zostały obliczone prawdopodobieństwa tego, że wartości zmiennej losowej $\bar{X} \sim N\left(5, 5; \frac{4}{\sqrt{n}}\right)$ w przy-

padku a) oraz zmiennej losowej $\bar{X} \sim N\left(6; \frac{4}{\sqrt{n}}\right)$ w przypadku b) przyjmą wartości z obszaru krytycznego – przedziału odrzucenia hipotezy zerowej (kolumna 2 w poniższej tablicy).

Tablica. Prawdopodobieństwa błędnej decyzji o odrzuceniu hipotezy zerowej w zależności od wielkości błędu systematycznego i liczebności próby

Liczebność próby (n)	Obszar krytyczny – przedział średniej z próby prowadzący do odrzucenia hipotezy zerowej	Prawdopodobieństwo błędnej decyzji o odrzuceniu hipotezy zerowej	
		błąd systematyczny = 0,5	błąd systematyczny = 1,0
16	$(-\infty; 3,04) \cup (6,96; +\infty)$	0,079	0,170
36	$(-\infty; 3,69) \cup (6,31; +\infty)$	0,117	0,323
80	$(-\infty; 4,12) \cup (5,88; +\infty)$	0,201	0,609
100	$(-\infty; 4,22) \cup (5,78; +\infty)$	0,240	0,705
200	$(-\infty; 4,45) \cup (5,55; +\infty)$	0,424	0,942
500	$(-\infty; 4,65) \cup (5,35; +\infty)$	0,798	≈ 1

Źródło: obliczenia własne.

W tym przykładzie jest widoczna wyraźna tendencja do wzrostu prawdopodobieństwa podjęcia błędnej decyzji o odrzuceniu testowanej hipotezy wraz ze wzrostem liczebności próby. Im mniejszy błąd systematyczny, tym ryzyko błędnej decyzji jest mniejsze.

Wspomniana wcześniej warunkowość prawidłowości statystycznych zaobserwowanych w próbie, odnosząca się do nieobciążenia obserwacji błędami systematycznymi, powinna być podkreślana każdorazowo przy interpretacji *p-value*. Wskaźnik ten wyraża prawdopodobieństwo uzyskania takiej próby, jaką zaobserwowano, albo bardziej ekstremalnej (bardziej nieprzystającej do hipotezy zerowej), pod warunkiem że nie tylko hipoteza zerowa jest prawdziwa, lecz także spełnione są wszystkie inne założenia w modelu wnioskowania¹⁴. W szczególności obejmują one losowe generowanie obserwacji w próbie oraz niewystępowanie błędów systematycznych w dokonanych pomiarach. Często jednak o tej drugiej części założeń zapomina się nie tylko przy interpretacji *p-value*, lecz także w całej procedurze wnioskowania, w której przez długi czas nie przywiązywano wagi do błędów nielosowych.

Kwestia wrażliwości testów statystycznych na błędy nielosowe w dużych próbach powoduje, że ta wrażliwość staje się jeszcze jednym argumentem na rzecz wzbogacenia testowania hipotez o inne elementy wnioskowania statystycznego. Przede wszystkim warto w tym kontekście wspomnieć o przedziałach ufności, które najczę-

¹⁴ Zwraca się na to uwagę m.in. w oświadczeniu Amerykańskiego Towarzystwa Statystycznego na temat istotności statystycznej oraz *p-value*, zob. Wasserstein i Lazar (2016, s. 131 i 132), a także Szreder (2019).

ściej dostarczają badaczowi więcej informacji niż tylko wskazanie, czy hipotezę dotyczącą danego parametru populacji należy odrzucić, czy uznać, że brak jest podstaw do jej odrzucenia. Jest to tym ważniejsze, że decyzja o braku podstaw do uznania hipotezy za fałszywą bywa niekiedy interpretowana implicite jako uznanie jej za prawdziwą. W celu zilustrowania tego, jakich informacji dostarcza badaczowi zastosowanie procedury weryfikacji hipotez, a jakich estymacja przedziałowa, rozważymy następujący przykład. Przyjmijmy, że testowana jest hipoteza o równości średnich w dwóch populacjach:

$$H_0: \mu_1 - \mu_2 = 0$$

$$H_1: \mu_1 - \mu_2 \neq 0$$

gdzie: μ_1 oznacza średnią arytmetyczną w pierwszej, a μ_2 – analogiczną średnią w drugiej populacji.

Założmy, że na podstawie wyników prób wylosowanych z tych populacji oszacowano 95-procentowy przedział ufności dla różnicy $\mu_1 - \mu_2$, który przybrał postać $[1, 5]$. Oznacza to, że badacz odrzuci sformułowaną wyżej hipotezę zerową na poziomie istotności 0,05, gdyż wartość 0 (ujęta w hipotezie zerowej) nie mieści się w przedziale liczbowym $[1, 5]$. Punktem centralnym tego przedziału, będącym próbkową oceną różnicy między średnimi w obu populacjach, jest wartość 3. Przyjmijmy następnie, że inna próba o tej samej liczebności dała tę samą wartość 3 jako ocenę różnicy między średnimi w obu populacjach, lecz rozproszenie wyników w próbie jest teraz większe, co powoduje, że przedział ufności ma postać $[-1, 7]$. Dla takiej próby zastosowanie procedury weryfikacji hipotez prowadzi do innej decyzji – mianowicie braku podstaw do odrzucenia hipotezy, że średnie w obu populacjach są identyczne. Jest tak dlatego, że wartość 0 mieści się w przedziale $[-1, 7]$. Czyli w tej drugiej sytuacji, mimo że różnica między średnimi próbkowymi wynosi 3 i jest taka, jak w pierwszym przypadku, badacz uzna, że nie ma podstaw do odrzucenia hipotezy zerowej. Dla niektórych odbiorców takiego komunikatu staje się to podstawą do uznania, że w dalszym procesie badawczym można założyć, że średnie te są równe. Tymczasem przedziały ufności znacznie lepiej obrazują to, co w przypadku weryfikacji hipotez kryje się za dychotomicznym wyborem – odrzucenia sprawdzanej hipotezy albo uznania, że brak jest podstaw do jej odrzucenia. W przypadku drugiej z rozważanych prób suchy komunikat z zastosowania testowania hipotez pozostawi wielu odbiorców w przekonaniu o dopuszczalności przyjęcia w przybliżeniu, że średnie w obu populacjach są identyczne. Natomiast przedział ufności sugeruje, że warto dokładnie to zbadać, ponieważ co prawda zerowa różnica między średnimi w populacjach nie jest – w świetle wyników próby – wykluczona, ale istnieje większe

prawdopodobieństwo, że różnica ta jest różna od 0 i dodatnia. Z tego i podobnych powodów coraz więcej czasopism naukowych oczekuje wzbogacenia testowania hipotez o przedstawienie przedziałów ufności, które pokazują pełny zakres różnic między tym, co zaobserwowano w próbie, a tym, co zapisano w hipotezie zerowej¹⁵. Ma to znaczenie także w sytuacjach, gdy dopuszcza się możliwość obciążenia danych w próbie błędami systematycznymi. Analiza przedziału ufności ogranicza ryzyko podjęcia błędnej decyzji opartej na procedurze testowania hipotez.

4. Podsumowanie

W empirycznych zastosowaniach wnioskowania statystycznego, w szczególności w badaniach ekonomicznych i społecznych, najbardziej uciążliwe są nie błędy losowe – których wielkość potrafimy wyrazić i zinterpretować – ale błędy nielosowe. Konsekwencje takich błędów, jak odmowy respondentów, braki w operacji losowania (błędy pokrycia), błędy treści czy powstałe na dalszych etapach badania błędy przetwarzania danych, stanowią współcześnie największe wyzwania dla statystyków. Wielu badaczy, świadomych tych błędów, usiłuje przekonać odbiorców wyników, że odpowiednio duża próba niweluje lub unieważnia konsekwencje ich oddziaływania. W artykule pokazano, że tak nie jest.

Testy statystyczne stosowane w dużych próbach losowych stają się bardziej niż w małych próbach wrażliwe zarówno na błędy losowe, jak i nielosowe. Badacz mający do dyspozycji dużą liczbę obserwacji w próbie bez trudu odrzuci hipotezę o nieistotności parametru populacji nawet wówczas, gdy w rzeczywistości jego wartość będzie bardzo bliska 0. Niewielka dyspersja statystyki testowej dla dużych prób sprawia bowiem, że dowolnie mała niezgodność próby z wynikami, jakie powinny się były pojawić, gdyby prawdziwa była hipoteza zerowa, doprowadzi do odrzucenia tej hipotezy. Jest to tylko kwestia odpowiedniej wielkości próby.

Wrażliwość ta odnosi się także do błędów o charakterze nielosowym, szczególnie błędów systematycznych. W dużych próbach każde obciążenie systematyczne w dokonanych obserwacjach skutkuje znacznym wzrostem ryzyka odrzucenia testowanej hipotezy w sytuacji, gdy w rzeczywistości jest ona prawdziwa. Ponadto im większa liczebność próby obciążonej błędem systematycznym, tym większe prawdopodobieństwo odrzucenia prawdziwej hipotezy zerowej. W artykule pokazano, że jednym ze sposobów zmniejszenia ryzyka błędnych decyzji w testowaniu hipotez może być analiza przedziału ufności przeprowadzona równolegle lub zamiast weryfikacji hipotez.

¹⁵ Amrhein, Trafimow i Greenland (2019, s. 266) postulują: „Interpretuj i zwracaj uwagę raczej na otrzymane oceny niż na testy, uważnie analizując wartości pomiędzy dolną i górną granicą przedziału ufności”.

Bibliografia

- Amrhein, V., Greenland, S., McShane, B. (2019). Retire statistical significance. *Nature*, (567), 305–307. <https://media.nature.com/original/magazine-assets/d41586-019-00857-9/d41586-019-00857-9.pdf>.
- Amrhein, V., Trafimow, D., Greenland, S. (2019). Inferential Statistics as Descriptive Statistics: There Is No Replication Crisis if We Don't Expect Replication. *The American Statistician*, 73(sup1), 262–270. <https://doi.org/10.1080/00031305.2018.1543137>.
- Bakan, D. (1966). The test of significance in psychological research. *Psychological Bulletin*, 66(6), 423–437. <https://doi.org/10.1037/h0020412>.
- Cohen, J. (1990). Things I have learned (so far). *American Psychologist*, 45(12), 1304–1312. <https://doi.org/10.1037/0003-066X.45.12.1304>.
- Dorofeev, S., Grant, P. (2006). *Statistics for Real-Life Sample Surveys. Non-Simple-Random Samples and Weighted Data*. Cambridge: Cambridge University Press. <https://doi.org/10.1017/CBO9780511543265>.
- Gelman, A., Stern, H. (2006). The Difference Between “Significant” and “Not Significant” is not Itself Statistically Significant. *The American Statistician*, 60(4), 328–331. <https://doi.org/10.1198/000313006X152649>.
- Gigerenzer, G., Marewski, J. N. (2015). Surrogate Science: The Idol of a Universal Method for Scientific Inference. *Journal of Management*, 41(2), 421–440. <https://doi.org/10.1177/0149206314547522>.
- Gigerenzer, G., Swijtink, Z., Porter, T., Daston, L., Beatty, J., Krüger, L. (1989). *The Empire of Chance: How probability changed science and everyday life*. Cambridge: Cambridge University Press. <https://doi.org/10.1017/CBO9780511720482>.
- Hirschauer, N., Grüner, S., Mußhoff, O., Becker, C., Jantsch, A. (2020). Can *p*-values be meaningfully interpreted without random sampling?. *Statistics Surveys*, 14, 71–91. <https://doi.org/10.1214/20-SS129>.
- Leamer, E. E. (1978). *Specification Searches: Ad Hoc Inference with Nonexperimental Data*. New York: John Wiley & Sons. https://www.anderson.ucla.edu/faculty_pages/edward.leamer/books/specification_searches.htm.
- Lempert, R. O. (2009). The Significance of Statistical Significance: Two Authors Restate An Incontrovertible Caution. Why a Book?. *Law & Social Inquiry*, 34(1), 225–249. <https://doi.org/10.1111/j.1747-4469.2009.01144.x>.
- Lin, M., Lucas, H. C., Jr., Shmueli, G. (2013). Too Big to Fail—Large Samples and the *p*-Value Problem. *Information Systems Research*, 24(4), 906–917. <https://doi.org/10.1287/isre.2013.0480>.
- Mercer, A. W., Kreuter, F., Keeter, S., Stuart, E. A. (2017). Theory and practice in nonprobability surveys. Parallels between causal inference and survey inference. *Public Opinion Quarterly*, 81(S1), 250–279. <https://doi.org/10.1093/poq/nfw060>.
- Nunnally, J. (1960). The place of statistics in psychology. *Educational and Psychological Measurement*, 20(4), 641–650. <https://doi.org/10.1177/001316446002000401>.
- Paradysz, J. (1989). O błędach nielosowych w badaniu dzietności kobiet w ramach Narodowego Spisu Powszechnego 1970. W: Główny Urząd Statystyczny, *Problemy badań statystycznych metodą reprezentacyjną* (s. 154–159). Warszawa.
- Paradysz, J. (2009). Błędy pokrycia w Narodowych Spisach Powszechnych. Statystyka w praktyce społeczno-gospodarczej. *Prace Naukowe / Akademia Ekonomiczna w Katowicach*, 65–76.

- Szreder, M. (2015). Zmiany w strukturze całkowitego błędu badania próbkowego. *Wiadomości Statystyczne*, 60(1), 4–12.
- Szreder, M. (2019). Istotność statystyczna w czasach big data. *Wiadomości Statystyczne. The Polish Statistician*, 64(11), 42–57. <https://ws.stat.gov.pl/Article/2019/11/042-057>.
- Vogt, W. P., Vogt, E. R., Gardner, D. C., Haeffele, L. M. (2014). *Selecting the Right Analyses for Your Data: Quantitative, Qualitative, and Mixed Methods*. New York: The Guilford Press.
- Wasserstein, R. L., Lazar, N. A. (2016). The ASA's Statement on p -Values: Context, Process, and Purpose. *The American Statistician*, 70(2), 129–133. <https://doi.org/10.1080/00031305.2016.1154108>.
- Ziliak, S. T., McCloskey, D. N. (2008). *The Cult of Statistical Significance: How the Standard Error Costs Us Jobs, Justice, and Lives*. Ann Arbor: University of Michigan. <https://doi.org/10.3998/mpub.186351>.

Czynniki determinujące dysproporcje w wynagrodzeniach kobiet i mężczyzn w krajach Unii Europejskiej

Dorota Witkowska^a, Aleksandra Matuszewska-Janica^b

Streszczenie. Głównym celem badania przedstawionego w artykule jest sprawdzenie, w jaki sposób wybrane czynniki determinujące nierówności implikowane płcią wpłynęły na rozmiar nieskorygowanej luki płacowej pracowników najemnych w Unii Europejskiej po kryzysie zapoczątkowanym w 2007 r. w odniesieniu do sytuacji sprzed kryzysu. Dodatkowym celem było wskazanie zmian w poziomie zatrudnienia kobiet i mężczyzn oraz w różnicach między płacami kobiet i mężczyzn, mierzonych wskaźnikiem *gender pay gap* (GPG), jakie nastąpiły w krajach UE po kryzysie w odniesieniu do okresu poprzedzającego kryzys. Badanie przeprowadzono za pomocą jednorównaniowych opisowych modeli ekonometrycznych określających lukę płacową. W analizie wykorzystano opublikowane przez Eurostat wyniki badań Structure of Earnings Survey (SES) i Labour Force Survey (LFS). Ze względu na dostępność danych przyjęto jako reprezentatywne dla sytuacji przed kryzysem dane za 2006 r. (uwzględniono kraje, które weszły do UE w późniejszych latach), a jako reprezentatywne dla sytuacji po kryzysie – dane za 2012 r. (wskaźnik zatrudnienia) oraz za 2012 r. i okres 2014–2018 (GPG).

Analizy wskaźników zatrudnienia kobiet i mężczyzn oraz różnic w płacach kobiet i mężczyzn wskazują, że po kryzysie nastąpił wzrost feminizacji zatrudnienia w 24 krajach UE, przy czym luka płacowa w latach 2006–2018 nie uległa istotnemu zmniejszeniu. Uzyskane wyniki pozwalają stwierdzić, że większa feminizacja zatrudnienia jest powiązana z większymi różnicami w płacach kobiet i mężczyzn. Podobna zależność występuje w przypadku współczynnika aktywizacji zawodowej. Dodatkowo obserwuje się istotne różnice w oddziaływaniu niektórych spośród rozpatrywanych czynników na rozmiar różnic płacowych w różnych grupach wieku i zawodowych.

Słowa kluczowe: rynek pracy, luka w płacach kobiet i mężczyzn, GPG, modele ekonometryczne, SES, LFS

JEL: C21, J31

Factors determining disproportions in men and women's wages in the European Union countries

Abstract. The primary aim of the presented study was to identify how selected factors determining gender-based inequalities affected the volume of the unadjusted pay gap among employees hired in the European Union after the 2007 crisis compared to the pre-crisis situation.

^a Uniwersytet Łódzki, Wydział Zarządzania, Katedra Zarządzania Finansami Przedsiębiorstwa / Univeristy of Lodz, Faculty of Management, Department of Entrepreneurship and Industrial Policy.
ORCID: <https://orcid.org/0000-0001-9538-9589>.

^b Szkoła Główna Gospodarstwa Wiejskiego w Warszawie, Instytut Finansów i Ekonomii, Katedra Ekonometrii i Statystyki / Warsaw University of Life Sciences – SGGW, Institute of Economics and Finance, Department of Econometrics and Statistics.

Autor korespondencyjny / Corresponding author, e-mail: aleksandra_matuszewska@sggw.edu.pl.
ORCID: <https://orcid.org/0000-0001-9127-6034>.

An additional purpose of the study was to indicate changes in the employment rates of men and women, as well as changes in the pay gap between the two sexes (measured by means of the gender pay gap index – GPG), which became noticeable in the EU countries after the crisis, as compared to the pre-crisis period. The study was conducted using single-equation descriptive econometric models describing the wage gap. The analysis was based on the results of the Structure of Earnings Survey (SES) and the Labour Force Survey (LFS), both published by Eurostat. Due to data availability issues, data for 2006 were assumed to be representative for the situation prior to the crisis (the study took into account also countries which became member states in later years), while data covering the year 2012 (employment rate) and the years 2014–2018 (GPG) were assumed as representative for the post-crisis period.

The analyses of the male and female employment rate and gender pay gaps indicate that following the crisis, the employment in the 24 EU countries became increasingly ‘feminised’, while no significant reduction of the pay gap was observed in the years 2006–2018. The obtained results indicate that greater ‘feminisation’ of employment is connected with greater gender pay gaps. A similar correlation occurs in relation to the professional activation rate. In addition, significant differences are observed in terms of the impact some of the analysed factors have on the volume of the gender wage gap in different age and occupational groups.

Keywords: labour market, gender pay gap, GPG, econometric models, SES, LFS

1. Wprowadzenie

Według szacunków Eurostatu w 2018 r. różnica między przeciętnymi płacami kobiet i mężczyzn w krajach Unii Europejskiej (UE-28), mierzona wartością wskaźnika *gender pay gap* (GPG), wynosiła 15%¹. Z danych o przeciętnych godzinowych wynagrodzeniach brutto kobiet i mężczyzn gromadzonych w ramach badania struktury wynagrodzeń (Structure of Earnings Survey – SES), na podstawie których oblicza się GPG², wynika, że poziom nierówności płacowych dla całego obszaru UE się zmniejszył. Zgodnie z obecnie dostępnymi danymi (Eurostat, 2021d, 2021e) w okresie między 2006 a 2018 r. spadek wartości wskaźnika GPG wyniósł ok. 2,7 p.proc. W literaturze przedmiotu (np. Klasen, 1999; Morrison i in., 2007) wskazuje się na istotne znaczenie zmniejszania się tych różnic dla gospodarki.

Rynek pracy, będący elementem systemu gospodarczego, jest niezwykle wrażliwy na koniunkturę gospodarczą. W okresie recesji obserwuje się wzrost bezrobocia, zmiany czasu pracy (częstsze są elastyczne formy pracy) czy spadek realnych wynagrodzeń (zob. m.in. Cazes i in., 2009; Zieliński, 2012), co może mieć wpływ na zróżnicowanie płac kobiet i mężczyzn.

Światowy kryzys gospodarczy, będący konsekwencją kryzysu finansowego zapoczątkowanego w 2007 r. w Stanach Zjednoczonych, miał daleko idące skutki dla

¹ Jest to nieskorygowana wartość luki płacowej – zob. wzór (1).

² W latach, w których to badanie zostało przeprowadzone, natomiast pomiędzy kolejnymi edycjami badania – na podstawie szacunkowych danych dostarczanych przez poszczególne kraje. Należy mieć na uwadze, że badanie SES nie obejmuje wszystkich pracujących w danej gospodarce, co zostanie wyjaśnione w dalszej części pracy.

krajów UE. W jego wyniku sektor finansowy odnotował straty, a gospodarki niemal wszystkich krajów członkowskich – spadek PKB. Znacząco zmalała produkcja, co pociągnęło za sobą likwidację wielu miejsc pracy oraz wzrost bezrobocia. W związku z tym rodzi się pytanie, czy w tych warunkach nastąpiła zmiana sytuacji kobiet na rynku pracy i jak kształtowała się ich pozycja po kryzysie. To zagadnienie jest przedmiotem badania omawianego w niniejszym artykule.

W ramach badania realizowano dwa cele. Głównym było sprawdzenie, w jaki sposób wybrane czynniki determinujące nierówności implikowane płcią wpłynęły na rozmiar nieskorygowanej luki płacowej pracowników najemnych w UE po kryzysie zapoczątkowanym w 2007 r. w odniesieniu do sytuacji sprzed kryzysu. Ze względu na dostępność danych okres badania obejmował lata 2006 i 2014. Dodatkowym celem było wskazanie zmian w poziomie zatrudnienia kobiet i mężczyzn oraz w różnicach między płacami kobiet i mężczyzn, mierzonych wskaźnikiem GPG, jakie nastąpiły w krajach UE po kryzysie w odniesieniu do okresu poprzedzającego kryzys. Ta część badania obejmowała lata 2006 i 2012. Analizę wskaźnika GPG uzupełniono o dane do 2018 r., aby wskazać tendencje w jego zmianach także w późniejszych latach.

Warto wspomnieć, że publikowane statystyki przeważnie wskazują na zmniejszanie się różnic na rynku pracy między obiema płciami. Jednak wiele dotychczasowych badań nie pokazuje polepszenia się pozycji kobiet. Na przykład Gago i Marcelo (2013) jako powód zmniejszenia się różnic wartości wskaźników zatrudnienia oraz bezrobocia kobiet i mężczyzn w Hiszpanii podają pogorszenie się sytuacji mężczyzn, przy niezmienionej pozycji kobiet. Podobne wnioski wyciąga Khitarishvili (2016) w przypadku Gruzji, stwierdzając, że za malejącą dysproporcję odpowiedzialne jest pogarszanie się sytuacji w sektorach zdominowanych przez mężczyzn. Z kolei Piazzalunga i Di Tommaso (2019), analizujący przyczyny zwiększenia się luki w płacach kobiet i mężczyzn we Włoszech w czasie kryzysu, upatrują je przede wszystkim w zamrożeniu płac w sektorze publicznym, który jest silnie sfeminizowany. Wymienione prace w głównej mierze opisują sytuację w pojedynczych krajach UE. Brakuje badań odnoszących się do tendencji na całym obszarze UE. Niniejsze opracowanie pozwoli częściowo wypełnić tę lukę.

2. Czynniki wpływające na różnice w płacach kobiet i mężczyzn w świetle badań literaturowych

Rozmiar nierówności między płacami kobiet i mężczyzn może się różnić w zależności od wielu czynników, takich jak: wykształcenie, wykonywany zawód, staż pracy, wymiar czasu pracy i wiek pracownika, a także wielkość przedsiębiorstwa i branża,

w jakiej działa. Staż pracy jest naturalnym czynnikiem różnicującym wynagrodzenia, ponieważ wraz z upływem czasu zdobywa się większe doświadczenie zawodowe, które powinno podwyższać wartość pracownika, a co za tym idzie – również jego wynagrodzenie. W przypadku kobiet staż pracy z reguły jest krótszy niż w przypadku mężczyzn. Dzieje się tak głównie ze względu na obowiązki opiekuńcze i wychowawcze, które przyczyniają się do całkowitej lub częściowej albo czasowej rezygnacji kobiet z pracy zawodowej (Anderson i in., 2001; Bukowski, 2010; Kotowska, 2007).

W literaturze naukowej prowadzona jest szeroka dyskusja, w której podejmowane są próby wyjaśnienia występowania na rynku pracy nierówności implikowanych płcią (np. Adamchik i Bedi, 2003; Becker, 1964; Blau, 2012; Blau i in., 2010; Blau i Kahn, 2006, 2017; Goraus i Tyrowicz, 2014; Grajek, 2003; Hirsch i in., 2010; Mortensen, 2012; Newell i Socha, 2007; Rubery i in., 1999). Dysproporcje w wynagrodzeniach kobiet i mężczyzn wyjaśnia się na gruncie wielu teorii, z których do najczęściej wykorzystywanych należą:

- teoria kapitału ludzkiego (Becker, 1962, 1964; Mincer, 1974; Schultz, 1961, 1971);
- teoria dyskryminacji (Becker, 1957);
- teoria preferencji (Charles i Grusky, 2004; Hakim, 2004, 2006; Jacobs i Gerson, 2004)³.

Badacze wskazują wiele czynników wpływających na dysproporcje w wynagrodzeniach kobiet i mężczyzn, które można zaklasyfikować do trzech grup:

- cechy indywidualne pracownika (np. wiek, staż pracy, typ i poziom wykształcenia, wykonywany zawód, czas poświęcany na pracę, status ekonomiczny i rodzinny, preferencje);
- cechy przedsiębiorstwa (np. rodzaj prowadzonej działalności, sektor publiczny/prywatny, wielkość przedsiębiorstwa, działalność związków zawodowych);
- cechy otoczenia (np. struktura zatrudnienia, sytuacja gospodarcza i społeczna w regionie/kraju, rozwiązania dotyczące instytucjonalnej opieki nad dziećmi i osobami starszymi, uwarunkowania kulturowe).

Różnorodność tych czynników powoduje, że dysproporcje mogą się różnić (nawet znacznie) w zależności od regionu czy grupy pracowników. Z badań wynika, że wśród młodszych pracowników dysproporcje te są mniejsze niż wśród pracowników starszych (Matuszewska-Janica, 2014), większe różnice obserwuje się wśród pracowników najlepiej zarabiających (Arulampalam i in., 2007), a w sektorze publicznym dysproporcje są mniejsze (Barón i Cobb-Clark, 2010). Duże zróżnicowanie występuje też w poszczególnych sektorach gospodarki (Matuszewska-Janica i Hozer-Koćmiel, 2015).

³ Teorie odnoszące się do kształtowania się wynagrodzeń zostały w sposób syntetyczny opisane w pracy Kryńskiej i Kopycińskiej (2015).

3. Metoda badania

Ogólnoświatowy kryzys, który wystąpił w pierwszej dekadzie XXI w., jest zjawiskiem szeroko opisywanym w literaturze naukowej (np. Lallement, 2011; Macdonald, 2012; Shostya, 2014, 2019; Terazi i Şenel, 2011). Za początek zawirowań gospodarczych zgodnie uznaje się rok 2007, mimo że pojawienie się ich negatywnych skutków w różnych krajach było rozłożone w czasie. Z tego powodu dla zobrazowania sytuacji przed kryzysem wybrano 2006 r. Z kolei za pierwszy rok po kryzysie przyjęto 2012 r., zgodnie z analizą zaprezentowaną w pracy Hozer-Koćmiel i in. (2017).

Jako miernik luki w płacach kobiet i mężczyzn posłużył wskaźnik GPG, który oblicza się zgodnie ze wzorem

$$GPG = (1 - \frac{GHE_F}{GHE_M}) \cdot 100, \quad (1)$$

gdzie GHE_F , GHE_M – przeciętne wynagrodzenie brutto za godzinę uzyskiwane odpowiednio przez kobiety i mężczyzn.

Jest to surowa (nieskorygowana) luka płacowa, ponieważ nie uwzględnia czynników, które różnicują poziom wynagrodzenia kobiet i mężczyzn (Blau i Kahn, 2007).

Badanie zostało podzielone na dwie części. W pierwszej analizowano zmiany w sytuacji kobiet na rynku pracy w latach 2006–2012 z uwzględnieniem takich czynników, jak: wskaźnik zatrudnienia kobiet i wskaźnik zatrudnienia mężczyzn wyznaczony dla grup wieku w przedziale 15–64 lat⁴ (zob. wzory (2a) i (2b)), udział zatrudnienia kobiet w głównych sektorach gospodarki (zob. wzór (3)) oraz rozmiar różnic w płacach kobiet i mężczyzn mierzonych wskaźnikiem GPG, z rozszerzeniem okresu badania do 2018 r. (zob. wzór (1)).

Wskaźnik zatrudnienia kobiet jest kalkulowany jako udział kobiet pracujących w ogóle kobiet w wieku 15–64 lat. Analogicznie wskaźnik zatrudnienia mężczyzn jest to udział pracujących mężczyzn w ogóle mężczyzn w wieku 15–64 lat:

$$E_{Fi} = \frac{NE_{Fi}}{NP_{Fi}} \cdot 100, \quad (2a)$$

$$E_{Mi} = \frac{NE_{Mi}}{NP_{Mi}} \cdot 100, \quad (2b)$$

⁴ Eurostat podaje ten wskaźnik dla różnych grup wieku. Na potrzeby prezentowanego badania przyjęto przedział 15–64 lat, ponieważ badanie Labour Source Survey (LFS), z którego zaczerpnięto dane do analizy, obejmuje osoby w tym przedziale wieku. W badaniu SES, którego wyniki również posłużyły do obliczeń, nie narzuca się kryterium wieku.

gdzie:

NE_{Fi} , NE_{Mi} – liczba zatrudnionych odpowiednio kobiet (F) i mężczyzn (M) w wieku 15–64 lat w i -tym kraju,

NP_{Fi} , NP_{Mi} – liczebność populacji odpowiednio kobiet (F) i mężczyzn (M) w wieku 15–64 lat w i -tym kraju.

Udział kobiet pracujących w poszczególnych sektorach gospodarki (dla przedziału wieku 15–64 lat) wyznaczono zgodnie ze wzorem

$$E\%_{Fij} = \frac{NE_{Fij}}{NE_{Fj}} \cdot 100, \quad (3)$$

gdzie:

NE_{Fij} – liczba kobiet w wieku 15–64 lat pracujących w i -tym kraju i j -ym sektorze gospodarczym,

NE_{Fi} – liczba kobiet w wieku 15–64 lat pracujących w i -tym kraju.

W tej części badania przyjęto podział na następujące sektory gospodarki: rolnictwo, przemysł i usługi. Wartości wskaźników zatrudnienia E_{Fi} i E_{Mi} oraz dane do obliczenia udziału kobiet pracujących w poszczególnych sektorach $E\%_{Fij}$ zaczerpnięto z badania LFS, które w Polsce jest realizowane pod nazwą Badania Aktywności Ekonomicznej Ludności. Dane te są publicznie dostępne w bazach danych Eurostatu⁵. Dane dotyczące wartości wskaźnika GPG pozyskano także z baz danych Eurostatu⁶, ale ze zbiorów odnoszących się do badania SES⁷. W odniesieniu do wskaźnika GPG podawanego przez Eurostat warto zaznaczyć, że nieskorygowany GPG jest obliczany przy wykorzystaniu danych dostarczanych przez SES co cztery lata, począwszy od 2006 r. Jest to punkt odniesienia dla szacunków krajowych (opartych na źródłach krajowych), które są dostarczane przez kraje członkowskie za lata między latami odniesienia SES (Eurostat, b.r.).

⁵ Wskaźniki zatrudnienia: Eurostat (2021a); liczba pracujących z podziałem na działy gospodarcze dla 2006 r.: Eurostat (2021b) i dla 2012 r.: Eurostat (2021c).

⁶ GPG (2006): Eurostat (2021d).

GPG (2009–2018): Eurostat (2021a).

⁷ W badaniu SES uwzględniane są firmy zatrudniające 10 osób lub więcej i dotyczy ono tylko pracowników najemnych, w odróżnieniu od badania LFS, w którym nie nałożono takich ograniczeń. Ponadto z badania SES (w przeciwieństwie do LFS) wyklucza się niektóre działy działalności gospodarczej, takie jak: rolnictwo; administracja publiczna, obrona narodowa, obowiązkowe ubezpieczenia społeczne; gospodarstwa domowe zatrudniające pracowników, produkujące wyroby i świadczące usługi na własne potrzeby; organizacje i zespoły eksterytorialne. Natomiast niewątpliwą zaletą badania SES jest pozyskiwanie danych z rejestrów przedsiębiorstw, co jest uznawane za bardziej wiarygodną informację niż ta pochodząca z badania LFS, w którym poziom wynagrodzenia jest deklarowany przez samych respondentów, czyli członków gospodarstw domowych z wylosowanych mieszkań, w wieku 15 lat i więcej.

W drugiej części badania analizowano wpływ wybranych czynników na poziom różnic w płacach kobiet i mężczyzn na obszarze UE w latach 2006 (przed kryzysem) i 2014 (po kryzysie)⁸, za pomocą opisowych jednorównaniowych modeli ekonometrycznych. Parametry modeli dla obu lat zostały oszacowane z wykorzystaniem metody najmniejszych kwadratów z korektą heteroskedastyczności na podstawie zagregowanych danych z badania SES pochodzących z 27 krajów UE⁹. Eurostat publikuje te dane w różnych układach, w odniesieniu do różnych cech, takich jak np.: sekcja działalności gospodarczej (według klasyfikacji NACE¹⁰), grupa wieku pracowników, grupa zaszerzegowania zawodowego (według klasyfikacji ISCO¹¹), poziom wykształcenia (według klasyfikacji ISCED¹²). Znaczna modyfikacja klasyfikacji działalności gospodarczej NACE¹³ w wersji Rev. 2 w stosunku do NACE Rev. 1.1, która nastąpiła w 2018 r., uniemożliwiła uwzględnienie w analizach danych dotyczących sektorów działalności gospodarczej przedsiębiorstw ze względu na brak porównywalności wyników uzyskanych dla lat 2006 i 2014. W związku z tym w analizie oprócz takich cech, jak:

- relacja liczby godzin przepracowanych (w miesiącu referencyjnym danego badania SES) przez kobiety do liczby godzin przepracowanych przez mężczyzn w wydzielonych grupach pracowników ($Activ_{ijk}$ – zob. wzór (8));
- stopień feminizacji wydzielonych grup pracowników (FEM_{ijk} – zob. wzór (7));
- poziom płac w wydrebnionych grupach pracowników w odniesieniu do średnich krajowych wynagrodzeń ($Wages_{ijk}$ – zob. wzór (6));

uwzględniono następujące charakterystyki:

- grupa wieku respondenta (Age_{ijk} – zob. tabl. 2);
- wykonywany zawód $ISCO_{ijk}$ (według grup zaszerzegowania zawodowego zgodnych z klasyfikacją ISCO – zob. tabl. 1).

W tak sporządzonych zbiorach danych obserwacje stanowiła grupa pracowników z i -tego kraju, j -ej grupy zawodowej i k -tej grupy wieku. Ostatecznie udało się wyodrębnić 1013 obserwacji dla 2006 r. (Eurostat, 2012f, 2021h, 2021j) oraz 1029 obserwacji dla 2014 r. (Eurostat, 2021g, 2021i, 2021k).

⁸ Wprawdzie w pracy Hozier-Koćmiel i in. (2017) za pierwszy rok po kryzysie uznano 2012 r., ale SES jest badaniem przeprowadzanym co cztery lata (począwszy od 2002 r.), zatem w badaniu można było uwzględnić dane dopiero z 2014 r.

⁹ W tej części badania pominięto Chorwację, która formalnie dołączyła do UE w 2013 r.

¹⁰ Nomenclature statistique des Activités économiques dans la Communauté Européenne – Statystyczna klasyfikacja działalności gospodarczych w Unii Europejskiej.

¹¹ The International Standard Classification of Occupations – Klasyfikacja zawodów i specjalności dla potrzeb rynku pracy. Zob. <https://www.ilo.org/public/english/bureau/stat/isco/>.

¹² The International Standard Classification of Education (ISCED) – Międzynarodowa Standardowa Klasyfikacja Edukacji. Zob. <https://unesdoc.unesco.org/ark:/48223/pf0000219109>.

¹³ Klasyfikacja NACE została szczegółowo przedstawiona w publikacji Eurostatu (2008). Zmiany, jakie nastąpiły w klasyfikacji NACE Rev. 2 w porównaniu z NACE Rev. 1.1 przedstawiono w rozdziale 5 (s. 45–50). Szczegółowe klasyfikacje NACE Rev. 1.1 i Rev. 2 można znaleźć także na stronach Eurostatu <https://ec.europa.eu/eurostat/data/classifications> w odnośniku RAMON, Eurostat's metadata server.

Postać modelu ekonometrycznego, według którego w badaniu szacowano parametry, jest następująca:

$$\ln WCR_{ijk} = \delta_0 + \delta_1 \ln Wages_{ijk} + \delta_2 \ln FEM_{ijk} + \delta_3 \ln Activ_{ijk} + \sum_{j=1}^8 \gamma_j ISCO_{ijk} + \sum_{k=1}^5 \lambda_k Age_{ijk} + \varepsilon_{ijk}, \quad (4)$$

gdzie:

$\ln WCR_{ijk}$ – logarytm naturalny współczynnika konwergencji płac kobiet i mężczyzn opisany wzorem (5a),

$\ln Wages_{ijk}$ – logarytm naturalny współczynnika poziomu wynagrodzeń opisany wzorem (6),

$\ln FEM_{ijk}$ – logarytm naturalny udziału kobiet w grupie pracowników opisany wzorem (7),

$\ln Activ_{ijk}$ – logarytm naturalny relacji liczby godzin przepracowanych przez kobiety w miesiącu referencyjnym do liczby godzin przepracowanych przez mężczyzn opisany wzorem (8),

$ISCO_{ijk}$ – zmienna binarna, która przyjmuje wartość 1, gdy współczynnik konwergencji płac kobiet i mężczyzn WCR_{ijk} odnosi się do j -ej grupy zawodowej (warianty tej zmiennej przedstawiono w tabl. 1),

Age_{ijk} – zmienna zero-jedynkowa, która przyjmuje wartość 1, gdy współczynnik konwergencji płac kobiet i mężczyzn WCR_{ijk} odnosi się do k -tej grupy wieku (warianty tej zmiennej przedstawiono w tabl. 2),

ε_{ijk} – składnik losowy modelu.

Logarytm naturalny współczynnika konwergencji płac kobiet i mężczyzn w i -tym kraju, j -ej grupie zawodowej i k -tej grupie wieku – $\ln WCR_{ijk}$ – oblicza się według wzoru

$$\ln WCR_{ijk} = \ln \frac{GHE_{F_{ijk}}}{GHE_{M_{ijk}}}, \quad (5a)$$

gdzie $GHE_{F_{ijk}}$, $GHE_{M_{ijk}}$ to przeciętne wynagrodzenie brutto za godzinę w i -tym kraju, w j -ej grupie zawodowej i w k -tej grupie wieku uzyskiwane odpowiednio przez kobiety i mężczyzn.

Wskaźnik WCR_{ijk} reprezentuje w modelu różnice w płacach kobiet i mężczyzn. Im bliższa 0 jest jego wartość, tym wyższa luka implikowana płcią na niekorzyść kobiet. Wprowadzenie wskaźnika w takiej postaci pozwoliło na przekształcenie war-

tości zmiennej z wykorzystaniem logarytmowania, ponieważ $WCR_{ijk} > 0$ dla każdego i, j , i k . W wyniku przekształcenia wskaźnika WCR_{ijk} otrzymuje się miernik nieskorygowanej luki w płacach kobiet i mężczyzn GPG:

$$GPG_{ijk} = (1 - WCR_{ijk}) \cdot 100. \quad (5b)$$

Logarytm naturalny współczynnika poziomu wynagrodzeń w i -tym kraju, j -ej grupie zawodowej i k -tej grupie wieku – $\ln Wages_{ijk}$ – służy do mierzenia poziomu płac w danej grupie zawodowej (zakłada się, że w grupach o przeciętnie wyższych zarobkach będą również większe różnice w wynagrodzeniach kobiet i mężczyzn). Jest on obliczany według wzoru

$$\ln Wages_{ijk} = \ln \frac{GHE_{ijk}}{GHE_i}, \quad (6)$$

gdzie:

GHE_{ijk} – przeciętne wynagrodzenie brutto za godzinę w i -tym kraju, j -ej grupie zawodowej i k -tej grupie wieku,

GHE_i – przeciętne wynagrodzenie brutto za godzinę w i -tym kraju.

Logarytm naturalny udziału kobiet w grupie pracowników z i -tego kraju, j -ej grupy zawodowej i k -tej grupy wieku – $\ln FEM_{ijk}$ – który jest wskaźnikiem określającym stopień feminizacji danej grupy pracowników, oblicza się zgodnie ze wzorem

$$\ln FEM_{ijk} = \ln \frac{EF_{ijk}}{EF_{ijk} + EM_{ijk}}, \quad (7)$$

gdzie EF_{ijk} , EM_{ijk} – liczba zatrudnionych w i -tym kraju, j -ej grupie zawodowej i k -tej grupie wieku odpowiednio kobiet i mężczyzn.

Z kolei jako rodzaj wskaźnika aktywizacji zawodowej można potraktować $\ln Activ_{ijk}$, za pomocą którego mierzy się relację przeciętnego czasu poświęcanego na pracę przez kobiety do czasu pracy mężczyzn:

$$\ln Activ_{ijk} = \ln \frac{ANHF_{ijk}}{ANHM_{ijk}}, \quad (8)$$

gdzie $ANHF_{ijk}$, $ANHM_{ijk}$ – czas pracy mierzony przeciętną liczbą godzin przepracowanych w miesiącu referencyjnym¹⁴ odpowiednio przez kobiety i mężczyzn zatrudnionych w i -tym kraju, j -ej grupie zawodowej i k -tej grupie wieku.

¹⁴ W badaniu SES czas pracy odnosi się do miesiąca referencyjnego, czyli tego, w którym zostało przeprowadzone badanie. W większości krajów jest to październik.

Tabl. 1. Główne grupy zawodowe według klasyfikacji ISCO-88

Warianty zmiennej	Grupy zawodów	Liczba obserwacji	
		2006	2014
ISCO 1	przedstawiciele władz publicznych, wyżsi urzędnicy i kierownicy	111	117
ISCO 2	specjaliści	108	130
ISCO 3	technicy i inni pracownicy średniego szczebla	129	124
ISCO 4	pracownicy biurowi	127	132
ISCO 5	pracownicy usług osobistych i sprzedawcy	127	128
ISCO 6	rolnicy, ogrodnicy, leśnicy i rybacy	56	43
ISCO 7	robotnicy przemysłowi i rzemieślnicy	121	112
ISCO 8	operatorzy i monterzy maszyn i urządzeń	113	115
ISCO 9	pracownicy przy pracach prostych (wariant referencyjny)	121	128
Suma		1013	1029

Uwaga. Uwzględniono liczbę obserwacji (rekordów) w próbie, które charakteryzowały się przynależnością do odpowiedniej grupy zaszerzgowania zawodowego według klasyfikacji ISCO.

Źródło: opracowanie własne na podstawie: Eurostat (2021f, 2021g, 2021h, 2021i, 2021j, 2021k).

Tabl. 2. Grupy wieku wyróżnione w badaniu SES

Warianty zmiennej	Grupy wieku w latach	Liczba obserwacji	
		2006	2014
Y0_29	< 30 (wariant referencyjny)	175	199
Y30_39	[30, 40)	222	218
Y40_49	[40, 50)	231	219
Y50_59	[50, 60)	217	215
Y60+	≥ 60	168	178
Suma		1013	1029

Uwaga. Uwzględniono liczbę obserwacji (rekordów) w próbie, które charakteryzowały się przynależnością do odpowiedniej grupy wieku.

Źródło: opracowanie własne na podstawie: Eurostat (2021f, 2021g, 2021h, 2021i, 2021j, 2021k).

Jak wspomniano, parametry modeli opisujących zróżnicowanie płac (4) zostały oszacowane dla danych za lata 2006 i 2014. W celu porównania parametrów modeli oszacowanych na danych z różnych lat wykorzystano test, w którym para hipotez przyjmuje postać (Aczel, 2000, s. 486; Tarczyński i in., 2013, s. 72):

$$H_0: \beta_i = \beta_{0i},$$

$$H_1: \beta_i \neq \beta_{0i}.$$

W wektorze parametrów β reprezentowane są parametry modelu (4), odpowiednio: $\delta_0, \delta_1, \delta_2, \delta_3, \gamma_j$ dla $j = 1, \dots, 8$ oraz λ_k dla $k = 1, \dots, 5$. W analizowanym przypadku przyjęto stałość informacji z 2006 r., zatem $\beta_{0i} = \beta_i^{2006}$ oraz $\beta_i = \beta_i^{2014}$, gdzie β_i^{2006} są parametrami modelu dla 2006 r., a β_i^{2014} – parametrami modelu dla 2014 r.

Sprawdzianem testu jest statystyka w postaci

$$t_i = \frac{b_i - b_{0i}}{S(b_i)} = \frac{b_i^{2014} - b_i^{2006}}{S(b_i^{2014})}, \quad (9)$$

gdzie:

b_i^{2006} , b_i^{2014} – parametry modeli szacowanych na danych odpowiednio z 2006 r. i 2014 r.,

$S(b_i^{2014})$ – standardowy błąd szacunku dla parametrów b_i^{2014} .

Weryfikację hipotez przeprowadzono dla poziomu istotności 0,05.

4. Zmiany wskaźników zatrudnienia i luki płacowej ze względu na płeć w krajach UE

Analizę aktywności kobiet i mężczyzn na rynku pracy w okresie uwzględniającym kryzys gospodarczy zapoczątkowany w 2007 r. przeprowadzono z wykorzystaniem wskaźnika zatrudnienia (zob. wzory (2a) i (2b), dotyczy on przedziału wieku 15–64 lat), a badania luki płacowej oparto na wskaźniku GPG. Odniesienia do danych, z których korzystano w tej części analizy, zostały przedstawione w rozdziale 3. Jak wcześniej podano, jako ostatni rok przed kryzysem przyjęto 2006, a jako pierwszy (dla UE) rok po kryzysie – 2012. Zmianę wartości wskaźnika zatrudnienia kobiet w tym okresie przedstawiono w tabl. 3. Zestawiono w niej także wartości wskaźnika zatrudnienia mężczyzn.

Tabl. 3. Zmiany wartości wskaźników zatrudnienia kobiet i mężczyzn w grupie osób w wieku 15–64 lat w okresie 2006–2012

K r a j e	Kobiety $K = E_{F_i}^{2012} - E_{F_i}^{2006}$	Mężczyźni $M = E_{M_i}^{2012} - E_{M_i}^{2006}$	Różnica ($K - M$)
	w p.proc.		
$K < 0$			
Grecja	-5,6	-13,8	8,2
Irlandia	-4,2	-15,2	11,0
Dania	-3,4	-6,0	2,6
Portugalia	-3,3	-9,2	5,9
Hiszpania	-2,6	-15,8	13,2
Słowenia	-1,3	-3,7	2,4
Chorwacja	-0,9	-3,5	2,6
Cypr	-0,9	-9,0	8,1
Estonia	-0,9	-1,7	0,8
Wielka Brytania	-0,9	-2,6	1,7
Rumunia	-0,2	3,0	-3,2
Łotwa	-0,1	-6,0	5,9

Tabl. 3. Zmiana wartości wskaźnika zatrudnienia kobiet w porównaniu ze zmianą wartości wskaźnika zatrudnienia mężczyzn w grupie osób w wieku 15–64 lat w okresie 2006–2012 (dok.)

Kraje	Kobiety $K = E_{\hat{F}_t}^{2012} - E_{\hat{F}_t}^{2006}$	Mężczyźni $M = E_{\hat{M}_t}^{2012} - E_{\hat{M}_t}^{2006}$	Różnica ($K - M$)
	w p.proc.		
$K > 0$			
Litwa	0,8	-4,2	5,0
Węgry	0,8	-2,3	3,1
Włochy	0,8	-4,1	4,9
Słowacja	0,8	-0,3	1,1
Finlandia	0,9	-0,9	1,8
Szwecja	1,1	0,1	1,0
UE-28 ^a	1,4	-1,9	3,3
Czechy	1,4	0,9	0,5
Francja	1,5	-0,9	2,4
Bułgaria	1,7	-1,5	3,2
Holandia	1,7	-1,6	3,3
Belgia	2,8	-1,0	3,8
Luksemburg	4,4	-0,1	4,5
Austria	4,5	1,3	3,2
Polska	4,9	5,4	-0,5
Niemcy	6,6	5,1	1,5
Malta	10,3	0,2	10,1

a W 2006 r. w skład UE wchodziło 25 krajów. Wartości wskaźników zatrudnienia dla agregatu UE-28 oszacowano przy uwzględnieniu informacji dla Bułgarii, Rumunii i Chorwacji.

Uwaga. Kraje uszeregowano według wartości zmiany wskaźnika dla kobiet.

Źródło: opracowanie własne na podstawie: Eurostat (2021a, 2021b, 2021c).

Można zauważyć, że po kryzysie – w 2012 r. – w UE-28, traktowanej jako agregat, wskaźnik zatrudnienia w analizowanej grupie wieku w przypadku kobiet wzrósł w stosunku do 2006 r. o 1,4 p.proc., a w przypadku mężczyzn spadł o 1,9 p.proc. We wszystkich krajach UE z wyjątkiem Polski i Rumunii różnice między zmianami wartości wskaźników zatrudnienia kobiet i mężczyzn były dodatnie. Oznacza to zaistnienie następujących sytuacji:

- zmniejszenie wartości wskaźnika zatrudnienia w przypadku kobiet było mniejsze niż w przypadku mężczyzn (11 krajów);
- zwiększenie wartości wskaźnika zatrudnienia w przypadku kobiet było większe niż w przypadku mężczyzn (5 krajów);
- wartość wskaźnika zatrudnienia kobiet się zwiększyła, a wartość wskaźnika zatrudnienia mężczyzn – zmniejszyła (10 krajów).

Największe różnice między wartościami wskaźnika zatrudnienia kobiet i mężczyzn zaobserwowano w Hiszpanii, która doświadczyła największego jego spadku w przypadku mężczyzn, i na Malcie, dla której wzrost tego wskaźnika był najwyższy w przypadku kobiet.

Warto odnotować, że odsetek pracujących kobiet w 2012 r. w stosunku do 2006 r. wzrósł w 16 krajach. W przypadku mężczyzn podobne zjawisko zaobserwowano tylko w sześciu krajach, a w dwóch innych w zasadzie nie nastąpiła zmiana. Biorąc pod uwagę odsetek kobiet pracujących w poszczególnych krajach UE, można zauważyć, że ich ekstremalne wartości w porównywanych latach uległy zmianie. Najmniejszą wartość wskaźnika zatrudnienia w 2006 r. odnotowano na Malcie (33,7%), a w 2012 r. – we Włoszech (41,7%). Z kolei najwyższą wartość wskaźnika zaobserwowano w krajach skandynawskich – 73,4% w Danii w 2006 r. i 71,8% w Szwecji w 2012 r. (w rozpatrywanych latach te kraje zamieniły się pozycjami). Najwięcej aktywnych zawodowo mężczyzn w 2006 r. było w Danii (81,2%) oraz w Holandii (80,9%), która w 2012 r. przejęła pierwszeństwo w rankingu, z zatrudnieniem na poziomie 79,3%, a zaraz po niej uplasowały się Niemcy, z wynikiem 77,9%. Przed kryzysem najniższym odsetkiem pracujących mężczyzn charakteryzowała się Polska (60,9%), a po kryzysie najniższą wartość wskaźnika zatrudnienia mężczyzn zaobserwowano w Chorwacji (58,5%), Grecji (60,1%) i Hiszpanii (60,3%).

Tabl. 4. Wartości wskaźnika zatrudnienia kobiet oraz odsetek kobiet pracujących w poszczególnych sektorach gospodarki w wybranych krajach UE

Kraje	Wskaźnik zatrudnienia kobiet		Udział kobiet pracujących w poszczególnych sektorach gospodarki w ogóle pracujących kobiet					
			rolnictwo		usługi		przemysł	
	2006	2012	2006	2012	2006	2012	2006	2012
	w %							
Dania	73,4	70,1	1,4	1,0	86,6	89,3	11,7	9,6
Grecja	47,3	41,7	12,3	12,6	77,7	79,8	10,0	7,5
Hiszpania	53,8	51,2	3,2	2,4	85,1	88,5	11,7	9,1
Irlandia	59,3	55,1	1,3	1,3	86,9	89,6	11,3	8,8
Malta	33,7	44,0	0,4	0,3	84,3	88,2	15,5	10,9
Niemcy	61,5	68,1	1,5	1,0	82,4	84,7	16,1	14,2
Polska	48,2	53,1	14,3	11,3	68,0	72,4	17,7	16,2
UE-28 ^a średnia	57,6	58,4	5,2	4,3	78,5	81,8	16,0	13,4
UE-28 ^a Vs	14,3	13,2	112,4	130,1	13,1	10,9	39,7	42,1

a W 2006 r. w skład UE wchodziło 25 krajów. Wartości wskaźników zatrudnienia dla agregatu UE-28 oszacowano przy uwzględnieniu informacji dla Bułgarii, Rumunii i Chorwacji.

Uwaga. Kraje uszeregowano alfabetycznie. Vs – współczynnik zmienności dla odchylenia standardowego. Obliczone średnia i odchylenie standardowe są miarami nieważonymi. Średnia wartość wskaźnika zatrudnienia (UE-28 średnia) w grupie mężczyzn wynosiła odpowiednio 71,5% w 2006 r. (przy Vs = 7,6%) oraz 69,6% w 2012 r. (przy Vs = 8,4%).

Źródło: opracowanie własne na podstawie: Eurostat (2021a, 2021b, 2021c).

W tabl. 4 przedstawiono wartości wskaźnika zatrudnienia kobiet w wybranych krajach UE oraz udział kobiet pracujących w poszczególnych sektorach gospodarki w ogóle pracujących kobiet. Do prezentacji wybrano te kraje, które wyróżniały się na

tle pozostałych. Szersza analiza tego problemu została przedstawiona w pracy Witkowskiej i in. (2019, s. 28–33). Zróżnicowanie wskaźnika zatrudnienia (w porównywanych latach, mierzone jako rozstęp) wśród kobiet było znacznie większe (30,1–39,7%) niż wśród mężczyzn (20,3–20,8%). Jednak w 2012 r. różnice występujące między krajami UE w przypadku kobiet zmniejszyły się, a w przypadku mężczyzn – nieznacznie wzrosły. Świadczą o tym również wartości współczynnika zmienności wyznaczone dla nieważonych średnich obliczonych dla UE – wartość wskaźnika dla kobiet zmalała z 14,3% do 13,2%, a dla mężczyzn wzrosła z 7,7% do 8,5%.

W odniesieniu do ogółu pracujących kobiet najniższy odsetek zaobserwowano w rolnictwie (średnia nieważona dla całej UE wyniosła 4,3–5,2%), przy czym zróżnicowanie tej cechy w próbie obejmującej analizowane kraje UE, mierzone współczynnikiem zmienności V_s dla odchylenia standardowego, było duże (ponad 100%). W przemyśle odsetek pracujących kobiet był nieco większy (13,4–16,0%) i odznaczał się średnim zróżnicowaniem (ok. 40%). Natomiast średni udział kobiet zatrudnionych w sektorze usług w ogóle pracujących kobiet wynosił aż 78,5% w 2006 r. i 81,8% w 2012 r. i charakteryzował się bardzo małym zróżnicowaniem (niewiele ponad 10%). Generalnie we wszystkich krajach UE zatrudnienie kobiet wzrosło w usługach, a zmniejszyło się w rolnictwie i przemyśle. Podobną tendencję zaobserwowano w Polsce.

Tabl. 5. Luka w płacach kobiet i mężczyzn w wybranych krajach UE

Kraje	2006	2009	2010	2012	2014	2015	2016	2017	2018
	w %								
Estonia	29,8	26,6	27,7	29,9	28,1	26,9	25,3	25,6	22,7
Czechy	23,4	25,9	21,6	22,5	22,5	22,5	21,5	21,1	20,1
Niemcy	22,7	22,6	22,3	22,7	22,3	22,0	21,5	21,0	20,9
Cypr	21,8	17,8	16,8	15,6	14,2	14,0	13,9	13,7	13,7
Grecja	20,7	.	15,0	.	12,5	.	.	.	.
Hiszpania	17,9	16,7	16,2	18,7	14,9	14,2	15,1	14,0	14,0
Dania	17,6	16,8	15,9	16,8	16,0	15,1	15,0	14,7	14,5
Irlandia	17,2	12,6	13,9	12,2	13,9	13,9	14,2	14,4	.
Luksemburg	10,7	9,2	8,7	7,0	5,4	5,5	5,5	5,0	4,6
Portugalia	8,4	10,0	12,8	15,0	14,9	17,8	17,5	16,3	16,2
Słowenia	8,0	–0,9	0,9	4,5	7,0	8,1	7,8	8,0	8,7
Polska	7,5	8,0	4,5	6,4	7,7	7,4	7,2	7,2	8,8
Malta	5,2	7,7	7,2	9,5	10,6	10,4	11,0	12,2	11,7
Włochy	4,4	5,5	5,3	6,5	6,1	5,5	5,3	5,0	.
Rumunia	7,8	7,4	8,8	6,9	4,5	5,8	5,2	3,5	3,0
UE-28	.	.	16,4	17,4	16,6	16,5	16,3	16,0	15,7
UE-27	17,7	.	16,5	17,3	16,7	16,4	.	.	.

Uwaga. Kraje uszeregowano według wartości wskaźnika GPG z 2006 r.

Źródło: opracowanie własne na podstawie: Eurostat (2021d, 2021e).

W tabl. 5 przedstawiono wartości luki płacowej (wskaźnika GPG) w UE oraz w wybranych krajach. W przypadku całej UE obserwuje się bardzo powolne zmniejszanie się luki płacowej, która w 2018 r. nadal była wysoka – 15,7%. Wśród krajów UE najniższą wartość GPG w analizowanym okresie odnotowano w Słowenii (–0,9% w 2009 r. i 0,9% w 2010 r.) oraz w Rumunii (3,5% w 2017 r. i 3,0% w 2018 r.). Ciekawym przypadkiem jest ujemna wartość wskaźnika GPG w Słowenii. Oznacza to, że w 2009 r. przeciętna płaca kobiet w tym kraju była wyższa od przeciętnej płacy mężczyzn o 0,9%. Z kolei najwyższe, ponad 20-procentowe, różnice w płacach kobiet i mężczyzn w całym analizowanym okresie odnotowano w Estonii, Czechach i Niemczech. Do 2010 r. do tej grupy należała również Finlandia, a do 2012 r. – Słowacja. Porównując wskaźniki GPG z lat 2012 i 2006, można zauważyć, że największy wzrost różnic w płacach kobiet i mężczyzn nastąpił w Portugalii (o 6,6 p.proc.) i na Malcie (o 4,3 p.proc.), a największy spadek tej różnicy – w Irlandii (o 5 p.proc.). W 2014 r. w porównaniu z 2006 r. najbardziej widoczne zwiększenie się luki płacowej odnotowano ponownie w Portugalii i na Malcie (odpowiednio o 6,5 i 5,4 p.proc.), a największe jej zmniejszenie się – w Luksemburgu (o 5,3 p.proc.). W okresie 2006–2018 największy wzrost wartości wskaźnika GPG (aż o 7,8 p.proc.) znowu dotyczył Portugalii, a największe spadki – Cypru (o 8,1 p.proc.), Estonii (o 7,1 p.proc.) i Luksemburga (o 6,1 p.proc.)¹⁵.

5. Wyniki estymacji parametrów modeli ekonometrycznych

W tabl. 6 przedstawiono wyniki estymacji parametrów modelu (1), który w dalszej części będzie określany jako model pełny¹⁶, dla prób zawierających wszystkie dostępne dane odpowiednio dla lat 2006 i 2014. Parametry przy zmiennych ilościowych ($\ln Wages_{ijk}$ – zob. wzór (6), $\ln FEM_{ijk}$ – zob. wzór (7), $\ln Activ_{ijk}$ – zob. wzór (8)) są ujemne i istotne. Oznacza to, że luka płacowa jest tym większa, im wyższe są wartości następujących wskaźników:

- przeciętnego wynagrodzenia brutto za godzinę w i -tym kraju, j -ej grupie zawodowej i k -tej grupie wieku w stosunku do średniej krajowej ($\ln Wages_{ijk}$);
- stopnia feminizacji, opisującego udział kobiet w grupie pracowników z i -tego kraju, j -ej grupy zawodowej i k -tej grupy wieku wśród wszystkich pracowników (obojsza płci) tak zdefiniowanej grupy ($\ln FEM_{ijk}$);

¹⁵ Bardziej szczegółowa analiza zróżnicowania luki płacowej w krajach UE została przedstawiona w pracy Witkowskiej i in. (2019, s. 121–152).

¹⁶ W modelach, których parametry szacowane są na próbach przekrojowych, uwzględniających znaczną liczbę zmiennych binarnych, istnieje podejrzenie występowania współliniowości. W związku z tym w omawianej analizie empirycznej wykorzystano współczynnik VIF do oceny występowania współliniowości. Nie potwierdził on obecności tego zjawiska.

- relacji liczby godzin przepracowanych przez kobiety do liczby godzin przepracowanych przez mężczyzn (w miesiącu referencyjnym danego badania SES) w wydzielonych grupach pracowników ($\ln Activ_{ijk}$).

Ponadto można zaobserwować, że wpływ stopnia feminizacji był silniejszy w 2014 r. w porównaniu z 2006 r., co wynika z istotnej różnicy między parametrami stojącymi przy tej zmiennej w modelach dla poszczególnych lat (wartość parametru w modelu opisującym sytuację w 2014 r. jest istotnie większa). We wszystkich grupach wieku zdiagnozowano większą lukę płacową niż w przypadku grupy referencyjnej (osoby w wieku 29 lat lub mniej). Największą różnicę wartości wskaźnika GPG w odniesieniu do zbioru obserwacji referencyjnych odnotowano w przypadku pracowników między 41. a 50. rokiem życia. Oznacza to, że sytuacja właśnie tych pracowników najbardziej przyczynia się do zwiększenia rozmiaru dysproporcji w płacach kobiet i mężczyzn. Warto też zwrócić uwagę na to, że w 2014 r. przynależność do grup wieku 30–39 lat oraz 40–49 lat miała znacznie mniejszą siłę oddziaływania na wartość wskaźnika GPG.

W analizie uwzględniającej podział na zawody zaobserwowano, że istotnie większe oddziaływanie na GPG niż w przypadku zmiennej referencyjnej (grupy pracowników zatrudnionych przy pracach prostych – ISCO 9) występuje w następujących grupach zawodowych:

- specjaliści (ISCO 2, zarówno w modelu opisującym sytuację w 2006 r., jak i w modelu odnoszącym się do 2014 r.);
- technicy i inni pracownicy średniego szczebla (ISCO 3, tylko w modelu dla 2006 r.);
- pracownicy biurowi (ISCO 4, w obu modelach);
- pracownicy usług osobistych i sprzedawcy (ISCO 5, w obu modelach);
- rolnicy, ogrodnicy, leśnicy i rybacy (ISCO 6, tylko w modelu dla 2006 r.).

Z kolei mniejsze oddziaływanie zdiagnozowano w grupach pracowników będących robotnikami przemysłowymi i rzemieślnikami (ISCO 7, w obu modelach) oraz operatorami i monterami maszyn i urządzeń (ISCO 8, w obu modelach).

Warto zauważyć, że oddziaływanie grupy wyższych urzędników i kierowników (ISCO 1), która wydaje się najlepiej uposażona, na rozmiar różnic w płacach kobiet i mężczyzn nie różni się istotnie od oddziaływania najmniej zarabiających z grupy pracowników zatrudnionych przy pracach prostych (ISCO 9). Może to wynikać ze znacznego zróżnicowania wynagrodzeń w obu grupach. Trzeba też dodać, że istotne zmiany w czasie odnotowano w przypadku grup pracowników usług osobistych i sprzedawców (ISCO 5) oraz rolników, ogrodników, leśników i rybaków (ISCO 6). W pierwszej z wymienionych grup nastąpił wzrost siły oddziaływania w porównaniu z grupą referencyjną (ISCO 9), a w drugiej – spadek.

Tabl. 6. Wyniki estymacji parametrów modeli pełnych

Zmienne	Model 2006	Model 2014	Test o równości parametrów (t_i)
Stała	-0,2022***	-0,2083***	-0,3117
<i>ln Wages</i>	-0,0877***	-0,1044***	-0,7771
<i>ln FEM</i>	-0,0334***	-0,0628***	-2,8243***
<i>ln Activ</i>	-0,4784***	-0,3908***	0,8841
Wiek: Y30_39	-0,0702***	-0,0369***	3,5791***
Y40_49	-0,0830***	-0,0628***	1,9779**
Y50_59	-0,0708***	-0,0575***	1,2342
Y60+	-0,0373***	-0,0467***	-0,8894
Zatrudnienie: ISCO 1	0,0327	0,0221	-0,4533
ISCO 2	0,0914***	0,0613***	-1,4825*
ISCO 3	0,0439**	0,0219	-1,4086*
ISCO 4	0,0914***	0,1086***	1,3106*
ISCO 5	0,0389**	0,0759***	3,3537***
ISCO 6	0,0989***	-0,0060	-3,9861***
ISCO 7	-0,1438***	-0,1587***	-0,7606
ISCO 8	-0,0554***	-0,0694***	-0,9295
R^2	0,3050	0,3188	

Uwaga. *ln Wages* – relacja płac w wyodrębnionych grupach pracowników w odniesieniu do średnich krajowych wynagrodzeń (zob. wzór (6)); *ln FEM* – stopień feminizacji wydzielonych grup pracowników (zob. wzór (7)); *ln Activ* – relacja liczby godzin przepracowanych przez kobiety do liczby godzin przepracowanych przez mężczyzn w wydzielonych grupach pracowników w miesiącu referencyjnym (zob. wzór (8)). Poziom istotności, na którym następuje odrzucenie hipotezy o braku istotności parametru: *** – 0,01; ** – 0,05; * – 0,1.

Źródło: obliczenia własne w programie Gretl na podstawie: Eurostat (2021f, 2021g, 2021h, 2021i, 2021j, 2021k).

Na kolejnym etapie analizy starano się odpowiedzieć na pytanie, czy wyróżnione w modelach zmienne oddziałują z jednakową siłą na GPG w różnych grupach wieku, przy czym badania ograniczono do 2014 r. (tabl. 7). Okazało się, że wskaźnik płac jest istotny i ujemny w grupach wieku do 29 lat włącznie oraz 30–39 lat, a dodatni w grupie 40–49 lat. Wśród młodszych pracowników wyższe niż przeciętne płace prowadzą zatem do wzrostu luki płacowej, a w grupie wieku 40–49 lat – do spadku wartości GPG. Dodatkowo stopień feminizacji ma wpływ na zwiększanie się GPG (jest istotnie ujemny) we wszystkich grupach wieku z wyjątkiem najmłodszej. Z kolei wskaźnik aktywności zawodowej jest istotnie ujemny jedynie dla grupy 40–49 lat, a istotnie dodatni – dla najstarszych pracowników. W starszych grupach wieku, tj. pracowników w wieku 40 lat i więcej, widoczne jest istotnie mniejsze oddziaływanie pierwszych trzech grup zawodowych (ISCO 1–ISCO 3), grup robotników przemysłowych i rzemieślników (ISCO 7) oraz operatorów i monterów (ISCO 8) niż klasy najniżej szeregowanych (pracowników przy pracach prostych – ISCO 9), a w grupach do lat 40 to oddziaływanie jest większe lub nieistotnie różne (wyjątek stanowi grupa robotników przemysłowych i rzemieślników w grupie wieku 30–39 lat).

Tabl. 7. Oceny estymatorów parametrów dla grup wieku w 2014 r.

Zmienne	Y0_29	Y30_39	Y40_49	Y50_59	Y60+
Stała	-0,1952***	-0,3137***	-0,1550***	-0,1554***	-0,1441***
<i>ln Wages</i>	-0,1843***	-0,1868***	0,1484***	0,0987*	0,0391
<i>ln FEM</i>	-0,0086	-0,0804***	-0,0533**	-0,0592**	-0,0515***
<i>ln Activ</i>	-0,6387*	-0,2395	-0,3861**	0,0328	0,3565***
Zatrudnienie: ISCO 1	0,0261	0,1423***	-0,2596***	-0,2189***	-0,1995***
ISCO 2	0,0867**	0,1592***	-0,1514***	-0,1093**	-0,1069***
ISCO 3	0,0565*	0,1097***	-0,1249***	-0,1261***	-0,1178***
ISCO 4	0,0873***	0,1665***	0,0309	0,0308	0,0305
ISCO 5	0,0581**	0,0850***	0,0290	0,0401	0,0626**
ISCO 6	-0,0064	0,0350	0,0196	-0,1028	-0,0284
ISCO 7	-0,0585	-0,1243***	-0,2154***	-0,2256***	-0,2398***
ISCO 8	-0,0361	-0,0499	-0,1129***	-0,1146***	-0,1140***
R^2	0,3857	0,2737	0,2704	0,3255	0,4701

Uwaga. Jak przy tabl. 6.

Źródło: obliczenia własne w programie Gretl na podstawie: Eurostat (2021g, 2021i, 2021k).

6. Podsumowanie

W ramach badania omawianego w artykule realizowano dwa cele. Pierwszym było wskazanie zmian w poziomie zatrudnienia kobiet oraz w różnicach między płacami kobiet i mężczyzn mierzonych wskaźnikiem GPG, jakie nastąpiły w krajach UE po kryzysie zapoczątkowanym w 2007 r. w odniesieniu do okresu poprzedzającego ten kryzys. Ta część analizy obejmowała lata 2006 i 2012, a w przypadku różnicy w płacach porównano wartości GPG do 2018 r. Drugim celem było sprawdzenie, czy zmienił się wpływ wybranych czynników na poziom nieskorygowanej luki płacowej pracowników najemnych w UE po kryzysie w odniesieniu do sytuacji sprzed kryzysu. W tym przypadku, ze względu na dostępność danych, okres analizy obejmował lata 2006 i 2014. Ponieważ korzystano z danych zagregowanych, obserwacja została zdefiniowana jako grupa pracowników w i -tym kraju, j -ej grupie zawodowej i k -tej grupie wieku. Grupa czynników determinujących rozmiar dysproporcji w płacach kobiet i mężczyzn, które uwzględniono w analizie empirycznej, obejmowała: poziom płac w wyodrębnionych grupach pracowników w odniesieniu do średnich krajowych wynagrodzeń; stopień feminizacji wydzielonych grup pracowników; relację liczby godzin przepracowanych przez kobiety do liczby godzin przepracowanych przez mężczyzn w wydzielonych grupach pracowników (w miesiącu referencyjnym danego badania SES); wiek (grupę wieku); wykonywany zawód (reprezentowany przez grupę zawodową według klasyfikacji ISCO).

Na podstawie przeprowadzonych analiz empirycznych sformułowano następujące wnioski. Po pierwsze po kryzysie finansowym, który rozpoczął się w 2007 r., a które-

go zakończenie nastąpiło w 2009 r., w UE wśród osób w wieku 15–64 lat zwiększył się odsetek pracujących kobiet, a z uwagi na to, że zmniejszył się odsetek pracujących mężczyzn, nastąpił wzrost feminizacji (tę zmianę odnotowano w 24 krajach UE), przy czym implikowana płcią luka płacowa (według danych z SES)¹⁷ nie uległa istotnemu zmniejszeniu. Można zatem przypuszczać, że zwiększona aktywność kobiet w okresie po kryzysie, czyli po 2009 r., może być spowodowana chęcią obniżenia kosztów pracy przez pracodawców. Kobiety są zatrudniane, ponieważ z reguły godzą się na niższe wynagrodzenia.

Po drugie uzyskane wyniki wskazują, że im większa feminizacja grupy pracowników, tym większe różnice w płacach, o czym świadczy niższy współczynnik konwergencji. Pogłębienie się tego zjawiska w 2014 r. może być skutkiem wzrostu feminizacji w gospodarkach UE. Podobnie jest w przypadku współczynnika aktywizacji zawodowej. Im bardziej czas przepracowany przez kobiety jest zbliżony do czasu przepracowanego przez mężczyzn (lub go przewyższa), tym większe są dysproporcje w płacach brutto za godzinę. Dodatkowo parametry modelu ekonometrycznego dla stosunku wynagrodzenia w danej grupie pracowników do średniego wynagrodzenia w danym kraju (jako zmiennej objaśniającej) charakteryzują się zazwyczaj wartościami ujemnymi. Tym samym należy się spodziewać, że im wyższe są wynagrodzenia w grupie pracowników, tym większe dysproporcje w zarobkach kobiet i mężczyzn.

Po trzecie wartości wszystkich parametrów dla wyróżnionych grup wieku (w modelu pełnym) są istotnie mniejsze od 0, co oznacza, że w porównaniu z referencyjną grupą wieku (do 29 lat) w grupach starszych pracowników różnice w płacach kobiet i mężczyzn są istotnie większe. Warto też zwrócić uwagę na różnice parametrów w modelach dotyczących lat 2006 i 2014. W przypadku dwóch najstarszych grup pracowników (w wieku 50–59 lat oraz 60 lat i więcej) wartość parametru nie zmieniła się istotnie. Różnica między wartościami wskaźnika GPG w tych grupach i w grupie najmłodszej utrzymała się w 2014 r. na podobnym poziomie co w 2006 r. Z kolei wartości parametrów dla grup wieku 30–39 lat oraz 40–49 lat istotnie się zmniejszyły w 2014 r., co oznacza, że różnice między lukami płacowymi w tych grupach a luką płacową w grupie najmłodszej także uległy zmniejszeniu. Taką sytuację można interpretować jako osłabienie siły oddziaływania zmiennej wiek na wielkość luki płacowej.

¹⁷ Należy mieć na uwadze, że o ile w badaniu LFS (z którego zaczerpnięto wartości wskaźnika aktywności zawodowej) są brani pod uwagę wszyscy pracujący, o tyle badanie SES (źródło informacji o GPG) ogranicza się do pracowników najemnych z przedsiębiorstw zatrudniających co najmniej 10 pracowników. Ponadto SES nie obejmuje niektórych sekcji działalności gospodarczej, takich jak np. rolnictwo czy administracja publiczna i obrona narodowa.

Po czwarte różnica w sile oddziaływania grup zawodowych na GPG istotnie się zwiększyła w przypadku pracowników usług osobistych i sprzedawców (ISCO 5). Ponadto nastąpiła istotna zmiana tego oddziaływania w przypadku rolników, ogrodników, leśników i rybaków (ISCO 6)¹⁸, polegająca na wyrównaniu siły oddziaływania tej grupy zawodowej i pracowników zatrudnionych przy pracach prostych (ISCO 9).

W modelach szacowanych dla poszczególnych grup wieku lub pojedynczych grup zawodowych widoczne są zmiany w sile i kierunku oddziaływania zmiennych zarówno ilościowych, jak i jakościowych. W grupach pracowników poniżej 50. roku życia dysproporcje w wynagrodzeniach kobiet i mężczyzn są większe i zwiększają się w kolejnych starszych grupach wieku, osiągając maksimum dla pracowników w wieku 40–49 lat. Natomiast w grupach wieku powyżej 49 lat różnice te zaczynają się zmniejszać. Może to wynikać z faktycznego wieku dezaktywizacji zawodowej kobiet, które zaczynają przechodzić na emeryturę po 50. roku życia. Spośród kobiet w wieku 60 lat i więcej pracują zaś tylko te, które chcą i dla których jest to opłacalne.

Bibliografia

- Aczel, A. D. (2000). *Statystyka w zarządzaniu* (tłum. K. Bech i in.). Warszawa: Wydawnictwo Naukowe PWN.
- Adamchik, V. A., Bedi, A. S. (2003). Gender pay differentials during the transition in Poland. *Economics of Transition*, 11(4), 697–726. <https://doi.org/10.1111/j.0967-0750.2003.00162.x>.
- Anderson, T., Forth, J., Metcalf, H., Kirby, S. (2001). *The gender pay gap: Final report to the Women and Equality Unit*. London: Cabinet Office Women and Equality Unit.
- Arulampalam, W., Booth, A. L., Bryan, M. L. (2007). Is there a glass ceiling over Europe? Exploring the gender pay gap across the wage distribution. *Industrial and Labor Relations Review*, 60(2), 163–186. <https://doi.org/10.1177/001979390706000201>.
- Barón, J. D., Cobb-Clark, D. A. (2010). Occupational Segregation and the Gender Wage Gap in Private-and Public-Sector Employment: A Distributional Analysis. *Economic Record*, 86(273), 227–246. <https://doi.org/10.1111/j.1475-4932.2009.00600.x>.
- Becker, G. S. (1957). *The Economics of Discrimination*. Chicago: The University of Chicago Press.
- Becker, G. S. (1962). Investment in Human Capital: A Theoretical Analysis. *Journal of Political Economy*, 70(5, Part 2), 9–49. <https://doi.org/10.1086/258724>.
- Becker, G. S. (1964). *Human Capital: A Theoretical and Empirical Analysis*. Chicago: The University of Chicago Press.
- Blau, F. D. (2012). *Gender, Inequality, and Wages*. Oxford: Oxford University Press.
- Blau, F. D., Ferber, M. A., Winkler, A. E. (2010). *The Economics of Women, Men, and Work* (6th edition). Pearson.

¹⁸ Warto zwrócić uwagę, że w tej części badania grupa zawodowa ISCO 6 – rolnicy, ogrodnicy, leśnicy i rybacy – nie odnosi się do sekcji działalności gospodarczej, jaką jest rolnictwo, leśnictwo i rybactwo (sekcja A). W badaniu SES ta sekcja jest pomijana. Obejmuje ona osoby pracujące w innych sekcjach, ale wykonujące zadania zdefiniowane w grupie zawodowej ISCO 6.

- Blau, F. D., Kahn, L. M. (2006). The U.S. Gender Pay Gap in the 1990s: Slowing Convergence. *Industrial and Labor Relations Review*, 60(1), 45–66. <https://doi.org/10.1177/001979390606000103>.
- Blau, F. D., Kahn, L. M. (2007). The Gender Pay Gap: Have Women Gone as Far as They Can?. *Academy of Management Perspectives*, 21(1), 7–23. <https://doi.org/10.5465/amp.2007.24286161>.
- Blau, F. D., Kahn, L. M. (2017). The Gender Wage Gap: Extent, Trends, and Explanations. *Journal of Economic Literature*, 55(3), 789–865. <https://doi.org/10.1257/jel.20160995>.
- Bukowski, M. (red.). (2010). *Zatrudnienie w Polsce 2008. Praca w cyklu życia*, Warszawa: Centrum Rozwoju Zasobów Ludzkich.
- Cazes, S., Verick, S., Heuer, C. (2009). *Labour market policies in times of crisis* (ILO Employment Working Paper No. 35). https://www.ilo.org/employment/Whatwedo/Publications/working-papers/WCMS_114973/lang--en/index.htm.
- Charles, M., Grusky, D. B. (2004). *Occupational Ghettos: the Worldwide Segregation of Women and Men*. Stanford: Stanford University Press.
- Eurostat. (b.r.). *Glossary: Gender pay gap (GPG)*. [https://ec.europa.eu/eurostat/statistics-explained/index.php/Glossary:Gender_pay_gap_\(GPG\)](https://ec.europa.eu/eurostat/statistics-explained/index.php/Glossary:Gender_pay_gap_(GPG)).
- Eurostat. (2008). *NACE Rev. 2. Statistical classification of economic activities in the European Community* (Eurostat Methodologies and Working Papers). <https://ec.europa.eu/eurostat/documents/3859598/5902521/KS-RA-07-015-EN.PDF>.
- Eurostat. (2021a). Employment rates by sex, age and citizenship (%) [zbiór danych]. https://appsso.eurostat.ec.europa.eu/nui/show.do?dataset=lfsa_ergaedn&lang=en.
- Eurostat. (2021b). Employment by sex, age and economic activity (1983–2008, NACE Rev. 1.1) (1000) [zbiór danych]. https://appsso.eurostat.ec.europa.eu/nui/show.do?dataset=lfsa_egana&lang=en.
- Eurostat. (2021c). Employment by sex, age and economic activity (from 2008 onwards, NACE Rev. 2) – 1000 [zbiór danych]. https://appsso.eurostat.ec.europa.eu/nui/show.do?dataset=lfsa_egan2&lang=en.
- Eurostat. (2021d). Gender pay gap in unadjusted form by NACE Rev. 1.1 activity – structure of earnings survey methodology [zbiór danych]. https://appsso.eurostat.ec.europa.eu/nui/show.do?dataset=earn_gr_gpg&lang=en.
- Eurostat. (2021e). Gender pay gap in unadjusted form by NACE Rev. 2 activity – structure of earnings survey methodology [zbiór danych]. https://appsso.eurostat.ec.europa.eu/nui/show.do?dataset=earn_gr_gpg2&lang=en.
- Eurostat. (2021f). Mean hourly earnings by sex, age, occupation – NACE Rev. 1.1, C-O excluding L [zbiór danych]. https://appsso.eurostat.ec.europa.eu/nui/show.do?dataset=earn_ses06_14&lang=en.
- Eurostat. (2021g). Mean hourly earnings by sex, age and occupation – NACE Rev. 2, B-S excluding O [zbiór danych]. https://appsso.eurostat.ec.europa.eu/nui/show.do?dataset=earn_ses14_14&lang=en.
- Eurostat. (2021h). Mean monthly hours paid by sex, age, occupation – NACE Rev. 1.1, C-O excluding L [zbiór danych]. https://appsso.eurostat.ec.europa.eu/nui/show.do?dataset=earn_ses06_35&lang=en.
- Eurostat. (2021i). Mean monthly hours paid by sex, age and occupation – NACE Rev. 2, B-S excluding O [zbiór danych]. https://appsso.eurostat.ec.europa.eu/nui/show.do?dataset=earn_ses14_35&lang=en.

- Eurostat. (2021j). Number of employees by sex, age, occupation – NACE Rev. 1.1, C-O excluding L [zbiór danych]. https://appsso.eurostat.ec.europa.eu/nui/show.do?dataset=earn_ses06_03&lang=en.
- Eurostat. (2021k). Number of employees by sex, age, occupation and size class – NACE Rev. 2, B-S excluding O [zbiór danych]. https://appsso.eurostat.ec.europa.eu/nui/show.do?dataset=earn_ses14_03&lang=en.
- Gago, E. G., Marcelo, S. K. (2013). Women, Gender Equality and the Economic Crisis in Spain. W: M. Karamessini, J. Rubery (Eds.), *Women and Austerity: The Economic Crisis and the Future for Gender Equality* (s. 228–247). Taylor & Francis Group.
- Goraus, K., Tyrowicz, J. (2014). *Gender Wage Gap in Poland – Can It Be Explained by Differences in Observable Characteristics?* (University of Warsaw Faculty of Economic Sciences Working Papers, No. 11/2014). https://www.wne.uw.edu.pl/files/5213/9636/9789/WNE_WP128_2014.pdf.
- Grajek, M. (2003). Gender Pay Gap in Poland. *Economic Change and Restructuring*, 36(1), 23–44.
- Hakim, C. (2004). *Key Issues in Women's work: Female Diversity and the Polarisation of Women's Employment* (2nd edition). New York: Routledge.
- Hakim, C. (2006). Women, careers, and work-life preferences. *British Journal of Guidance & Counselling*, 34(3), 279–294. <https://doi.org/10.1080/03069880600769118>.
- Hirsch, B., Schank, T., Schnabel, C. (2010). Differences in Labor Supply to Monopsonistic Firms and the Gender Pay Gap: An Empirical Analysis Using Linked Employer-Employee Data from Germany. *Journal of Labor Economics*, 28(2), 291–330. <https://doi.org/10.1086/651208>.
- Hozer-Koćmiel, M., Halik, P., Sobolewska, A. (2017). Has the economic crisis equally influenced women's and men's employment in Poland on the background of the EU countries? W: S. Misiak-Kwit, M. Wiścicka-Fernando (Eds.), *Equality and Management* (s. 7–18). Szczecin: Volumina.pl.
- Jacobs, J. A., Gerson, K. (2004). *The Time Divide: Work, Family, and Gender Equality*. Cambridge: Harvard University Press.
- Khitarrishvili, T. (2016). Two tales of contraction: gender wage gap in Georgia before and after the 2008 crisis. *IZA Journal of Labor & Development*, 5, 1–28. <https://doi.org/10.1186/s40175-016-0060-z>.
- Klasen, S. (1999). *Does Gender Inequality Reduce Growth and Development? Evidence from Cross-Country Regressions* (World Bank Working Paper Series No. 7). <http://documents1.worldbank.org/curated/en/612001468741378860/pdf/multi-page.pdf>.
- Kotowska, I. E. (2007). Uwagi o polityce rodzinnej w Polsce w kontekście wzrostu dzietności i zatrudnienia kobiet. *Polityka Społeczna*, (8), 13–19.
- Kotowska, I. E. (red.). (2009). *Strukturalne i kulturowe uwarunkowania aktywności zawodowej kobiet w Polsce*. Warszawa: Wydawnictwo Naukowe Scholar.
- Kryńska, E., Kopycińska, D. (2015). Wages in Labour Market Theories. *Folia Oeconomica Stetinensia*, 15(2), 177–190. <https://doi.org/10.1515/fofi-2015-0044>.
- Lallement, M. (2011). Europe and the economic crisis: forms of labour market adjustment and varieties of capitalism. *Work, Employment and Society*, 25(4), 627–641. <https://doi.org/10.1177/0950017011419717>.
- Macdonald, R. (2012). *Genesis of the Financial Crisis*. Palgrave Macmillan.

- Matuszewska-Janica, A. (2014). Wages inequalities between men and women: Eurostat SES meta-data analysis applying econometric models. *Metody Ilościowe w Badaniach Ekonomicznych*, 15(1), 113–124. http://qme.sggw.pl/pdf/MIBE_T15_z1_11.pdf.
- Matuszewska-Janica, A., Hozer-Koćmiel, M. (2015). Struktura zatrudnienia oraz wynagrodzenia kobiet i mężczyzn a przedmiotowa struktura gospodarcza w państwach UE. *Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu*, (385), 178–186. <https://doi.org/10.15611/pn.2015.385.19>.
- Mincer, J. A. (1974). The human capital earnings function. W: J. A. Mincer, *Schooling, Experience, and Earnings* (s. 83–96). National Bureau of Economic Research. <https://www.nber.org/system/files/chapters/c1767/c1767.pdf>.
- Morrison, A., Raju, D., Sinha, N. (2007). *Gender Equality, Poverty and Economic Growth* (Policy Research Working Paper No. 4349). <https://openknowledge.worldbank.org/handle/10986/7321>.
- Mortensen, D. T. (2012). *Dyspersja płac: dlaczego podobni pracownicy zarabiają różnie?* (tłum. Z. Matkowski). Warszawa: Polskie Towarzystwo Ekonomiczne.
- Newell, A., Socha, M. (2007). *The Polish Wage Inequality Explosion* (IZA Discussion Papers No. 2644). <http://ftp.iza.org/dp2644.pdf>.
- Piazzalunga, D., Di Tommaso, M. L. (2019). The increase of the gender wage gap in Italy during the 2008–2012 economic crisis. *The Journal of Economic Inequality*, 17(2), 171–193. <https://doi.org/10.1007/s10888-018-9396-8>.
- Rubery, J., Smith, M., Fagan, C. (1999). *Women's Employment in Europe: Trends and Prospects*. London: Routledge.
- Schultz, T. W. (1961). Investment in human capital. *The American Economic Review*, 51(1), 1–17.
- Schultz, T. W. (1971). *Investment in Human Capital. The Role of Education and of Research*. New York: The Free Press.
- Shostya, A. (2014). The effect of the global financial crisis on transition economies. *Atlantic Economic Journal*, 42(3), 317–332. <https://doi.org/10.1007/s11293-014-9418-2>.
- Shostya, A. (2019). The Global Financial Crisis in Transition Economies: The Role of Initial Conditions. *Atlantic Economic Journal*, 47(1), 37–51. <https://doi.org/10.1007/s11293-019-09607-8>.
- Tarczyński, W., Witkowska, D., Kompa, K. (2013). *Współczynnik beta. Teoria i praktyka*. Pielaszek Research.
- Terazi, E., Şenel, S. (2011). The effects of the global financial crisis on the Central and Eastern European Union countries. *International Journal of Business and Social Science*, 2(17), 186–192. http://ijbssnet.com/view.php?u=https://ijbssnet.com/journals/Vol_2_No_17/25.pdf.
- Witkowska, D., Kompa, K., Matuszewska-Janica, A. (2019). *Sytuacja kobiet na rynku pracy: wybrane aspekty*. Łódź: Wydawnictwo Uniwersytetu Łódzkiego.
- Zieliński, M. (2012). Kryzys gospodarczy a sytuacja na rynku pracy w krajach Unii Europejskiej w latach 2007–2011. *Oeconomia Copernicana*, 3(2), 57–68. <http://dx.doi.org/10.12775/OeC.2012.009>.

Wykorzystanie Google Trends do modelowania stopy bezrobocia rejestrowanego w Polsce

Mariusz Malinowski^a

Streszczenie. W artykule podjęto problematykę monitorowania stopy bezrobocia w Polsce. Celem przedstawionego badania jest sprawdzenie, czy dołączenie wybranych indeksów Google Trends do autoregresyjnego modelu stopy bezrobocia rejestrowanego poprawia trafność generowanych przez niego prognoz. Zastosowana metoda badania opiera się na technikach nowcastingu służących do oceny bieżącego stanu gospodarki. Dane za lata 2004–2019 zostały zaczerpnięte z publikacji GUS oraz serwisu Google Trends, który pozwala na śledzenie popularności terminów wyszukiwanych przez internautów. Porównano jakość dopasowania modelu do danych oraz błędy prognoz modelu podstawowego i modeli rozszerzonych o zmienne egzogeniczne. Artykuł przedstawia zarówno potencjał, jak i ograniczenia wykorzystywania nowego źródła danych w analizach makroekonomicznych dotyczących Polski. Na podstawie przeprowadzonej analizy można uznać, że indeksy Google, powszechnie wykorzystywane w literaturze anglojęzycznej, nie poprawiają trafności predykcji modelu autoregresyjnego. Zadowalające rezultaty uzyskiwane są tylko dla indeksów związanych z międzynarodową mobilnością siły roboczej.

Słowa kluczowe: nowcasting, Google Trends, bezrobocie, dane makroekonomiczne

JEL: C22, C55, C82, E27

Use of Google Trends in modelling registered unemployment rate in Poland

Abstract. The paper deals with the problem of monitoring the unemployment rate in Poland. The main aim of the article is to check whether the addition of selected Google Trends indices improves the accuracy of forecasts generated by the autoregressive model of registered unemployment rate. The research method is based on nowcasting techniques which are used to assess the current state of an economy. Data for the years 2004–2019 were retrieved from publication by Statistics Poland (GUS) and Google Trends, the latter of which allows tracking the popularity of terms searched by Internet users. The study compares the goodness of fit and forecast errors of the basic model with these of models extended with exogenous variables. Both the potential and the limitations of the utilisation of a new source of data in macroeconomic analyses concerning Poland are presented in the paper. The analysis yields a conclusion that Google indexes (commonly used in the literature written in English) do not improve the accuracy of predictions of the autoregressive model. Satisfactory results are only obtained for indices related to the international mobility of the workforce.

Keywords: nowcasting, Google Trends, unemployment, macroeconomic data

^a Uniwersytet Warszawski, Wydział Nauk Ekonomicznych / University of Warsaw, Faculty of Economic Sciences. E-mail: mariusz.a.malinowski@uw.edu.pl. ORCID: <https://orcid.org/0000-0003-4635-2960>.

1. Wprowadzenie

Cyfrowe ślady pozostawiane przez każdego użytkownika internetu spowodowały w ostatnich latach wykładniczy wzrost liczby źródeł danych, które mogą być wykorzystywane w analizach ekonomicznych (Blazquez i Domenech, 2018). Badacze życia gospodarczego mogą w swoich pracach wyjść poza konwencjonalne ankiety czy oficjalne statystyki, opracowywane i publikowane głównie przez instytucje państwowe. Aktywność ogromnej liczby indywidualnych użytkowników, firm oraz organizacji publicznych może dostarczać danych pomagających opisać ich zachowania, decyzje i motywacje, a tym samym pomóc monitorować kluczowe ekonomiczne czy społeczne zmiany i trendy.

Pod koniec lat 90. XX w. ogrom danych generowanych przez użytkowników internetu za pośrednictwem ciągle rozwijającej się technologii cyfrowej zaczął być określany terminem *big data* (Cox i Ellsworth, 1997). Na początku XXI w. termin ten zdefiniowano w kategoriach modelu 3V (Laney, 2001), na który składają się: liczba danych (*volume*), prędkość przetwarzania danych (*velocity*) oraz różnorodność danych (*variety*). W następnych latach, wraz z rozwojem analityki dużych zbiorów danych, model rozszerzono do 4V, dodając wymiar wartości (*value*). Współcześnie koncepcja *big data* zaczyna być definiowana w kategoriach modelu 5V (Bello-Orgaz i in., 2016). Kolejną jego składową jest wiarygodność (*veracity*), która odnosi się do odpowiedniego zarządzania zbiorami danych oraz polityki prywatności.

Jedno ze źródeł *big data*, które są często wykorzystywane w kontekście badań ekonomicznych, stanowią dane dotyczące słów kluczowych wyszukiwanych w sieci. Użytkownik internetu wpisuje interesujące go słowo bądź frazę do wyszukiwarki, która następnie dostarcza mu informacji możliwie najlepiej dopasowanych do danych wejściowych. Firma Google w 2000 r. zrewolucjonizowała ten proces oraz ukształtowała sposób, w jaki dzisiaj przeszukujemy sieć w celu uzyskania konkretnych informacji (Buono i in., 2017). Google wprowadziło bowiem algorytm nazywany PageRank, który podpowiada użytkownikom strony internetowe na podstawie liczby odniesień do nich występujących na innych stronach.

Od maja 2006 r. dostępna jest usługa Google Trends, która dostarcza aktualnych raportów o wolumenie zapytań o dane słowa kluczowe i frazy w różnych regionach świata oraz różnych językach, z historią wyszukiwania dostępną od stycznia 2004 r. Serwis pozwala na generowanie szeregów czasowych obrazujących zmiany zainteresowania danym tematem. Użytkownik określa słowo kluczowe bądź frazę, a Google Trends pokazuje wykres liniowy z czasem na osi poziomej oraz częstotliwością wyszukiwania na osi pionowej.

Dane z Google Trends aktualizowane są w czasie rzeczywistym, w związku z czym często wykorzystuje się je w literaturze makroekonomicznej do monitorowania bie-

żącego stanu gospodarki, które w publikacjach anglojęzycznych określane jest terminem *nowcasting* (Giannone i in., 2008). Termin ten stanowi połączenie słów *now* oraz *forecasting* i w ostatnich latach coraz częściej pojawia się w artykułach ekonomicznych, co wynika ze wzrostu zapotrzebowania na dokładne krótkoterminowe analizy i prognozy stanu gospodarki (Kapetanios i Papailias, 2018). Kluczowe miary, takie jak PKB i jego składowe, inflacja czy stopa bezrobocia, publikowane są z dość znacznym opóźnieniem, a często – po jeszcze dłuższym czasie – poddawane są rewizji. Dlatego coraz chętniej sięga się po nietypowe, ale łatwo dostępne, aktualne i wiarygodne źródła przydatne do wstępnych estymacji.

Celem badania omawianego w artykule jest sprawdzenie, czy dołączenie wybranych indeksów Google Trends do autoregresyjnego modelu stopy bezrobocia rejestrowanego poprawia trafność generowanych przez niego prognoz. W badaniu analizowana jest użyteczność indeksu praca oraz indeksów alternatywnych, wcześniej niewykorzystywanych w literaturze naukowej dotyczącej prognozowania stopy bezrobocia. Termin *praca*, z którego korzystają np. Pavlicek i Kristoufek (2015), jest bezpośrednim tłumaczeniem terminu *jobs*, najpowszechniej używanego w anglojęzycznych publikacjach poruszających tę tematykę. Warto jednak zauważyć, że *praca* w języku polskim jest stosunkowo dość często wykorzystywana w kontekście innym niż działalność zarobkowa. W języku angielskim występuje rozróżnienie na *job* i *work*, dzięki któremu wyszukiwane przez uczniów słowo kluczowe *homework* nie ma wpływu na zmienność szeregu używanego w badaniu. Natomiast w języku polskim można mówić nie tylko o pracy domowej, lecz także licencjackiej czy magisterskiej. Jest to też określenie funkcjonowania organizmu czy urządzenia, a nawet wielkość fizyczna określająca ilość energii potrzebną do przemieszczenia się obiektu.

Należy również zwrócić uwagę na charakterystyki polskiej siły roboczej, o których wspominają Pavlicek i Kristoufek (2015). Nie ulega wątpliwości, że po akcesji Polski do Unii Europejskiej w 2004 r. nasiliły się migracje mieszkańców Polski oraz zmieniła się struktura społeczno-demograficzna tych migracji. Jak wynika z *Informacji o rozmiarach i kierunkach czasowej emigracji z Polski w latach 2004–2017* (Główny Urząd Statystyczny [GUS], 2018), liczba osób przebywających czasowo za granicą zwiększyła się z 1 mln w 2004 r. do 2 mln w 2010 r. i rosła aż do 2,54 mln w 2017 r. Bartosik (2012) zwraca uwagę, że w przypadku Polski związek rynku pracy i mobilności międzynarodowej jest dwojakiego rodzaju. Z jednej strony zła sytuacja na rynku pracy (tj. wysokie bezrobocie) może stanowić czynnik wypychający (innymi słowy, sprzyjający wyjazdowi), a z drugiej – międzynarodowa mobilność siły roboczej wpływa na zmniejszanie się podaży pracy, co z kolei sprzyja spadkowi bezrobocia oraz wzrostowi płac.

W związku ze wspomnianymi problemami z terminem *praca* oraz charakterystyczną dla Polski międzynarodową mobilnością siły roboczej w badaniu wykorzy-

stywany jest również szereg czasowy *praca za granicą*. Na potrzeby artykułu przyjęto założenie, że bezrobotni najpierw szukają ogólnych opinii o możliwościach zarobkowania poza granicami kraju i najczęściej wpisują najprostsze określenie: *praca*.

Problematiczność terminu *praca* skłania do wyeliminowania tego słowa kluczowego i zastąpienia go alternatywnym – mniej wieloznacznym – terminem, który jednak dość dobrze wpisuje się w proces poszukiwania zatrudnienia. W związku z tym w niniejszym badaniu wykorzystywano szereg CV + *curriculum vitae*, przy założeniu, że najczęstszym pierwszym krokiem podejmowanym przez bezrobotnych podczas poszukiwania zatrudnienia jest właśnie przygotowanie tego dokumentu.

2. Przegląd literatury

Jeszcze zanim Google Trends stało się publicznie dostępnym serwisem, Ettredge i in. (2005) wykorzystali dane z *WordTracker's Top 500 Keyword Report*, opublikowanego przez Rivergold Associates, do analizy ekonometrycznej. Jednak ze względu na ograniczoną liczbę obserwacji w ich pracy nie można znaleźć analizy szeregów czasowych ukierunkowanej na predykcje. Autorzy badają siłę objaśniającą modeli, a jednocześnie wykazują, że istnieje pozytywna i statystycznie istotna korelacja pomiędzy słowami kluczowymi dotyczącymi poszukiwania pracy, które indywidualni internauci ze Stanów Zjednoczonych wpisują w wyszukiwarki, a oficjalnymi statystykami dotyczącymi bezrobocia.

Korzystanie z Google Trends w analizach ekonometrycznych zapoczątkowali Choi i Varian (2009), którzy badali użyteczność tego nowego źródła danych w prognozowaniu poziomów sprzedaży detalicznej, samochodów i domów w Stanach Zjednoczonych oraz liczby turystów w Hongkongu.

W tym samym roku Askitas i Zimmermann (2009) przebadali korelację pomiędzy różnymi wyszukiwanymi słowami kluczowymi a stopą bezrobocia w Niemczech w okresie od stycznia 2004 r. do kwietnia 2009 r. Wykazują oni w swoim artykule, że indeksy Google Trends z trzeciego i czwartego tygodnia poprzedniego miesiąca mogą być wykorzystywane do prognozowania stopy bezrobocia w bieżącym miesiącu. Co warto podkreślić, ta praca była pierwszą, w której zestawiono ze sobą statystyki dotyczące stopy bezrobocia i dane pochodzące z Google Trends.

D'Amuri i Marcucci (2012) przeprowadzili szczegółowe badanie siły predykcyjnej ponad 500 różnych liniowych i nieliniowych modeli prognozujących stopę bezrobocia w Stanach Zjednoczonych, przy czym w większości z nich do zbioru zmiennych objaśniających włączyli szeregi czasowe Google Trends. Ich praca jest szczególnie ważna w kontekście metody badania przedstawionej w niniejszym artykule, ponieważ wskazuje główne kroki na etapie analizowania użyteczności zmiennych Google. Autorzy przeprowadzili swoje badanie na podstawie następującego schematu: dobór

słów kluczowych (tj. wygenerowanie odpowiednich szeregów czasowych Google Trends), konstrukcja i dopasowanie modeli oraz ostateczne porównanie ich siły predykcyjnej w prognozie *out-of-sample*. Udowodnili, że modele wykorzystujące dane Google charakteryzują się większą dokładnością w przewidywaniu poziomu stopy bezrobocia.

Pavlicek i Kristoufek (2015) wykorzystali wspomniany schemat metodologiczny do analizy zasadności wykorzystywania danych Google w nowcastingu stopy bezrobocia w państwach Grupy Wyszehradzkiej. W swojej pracy dowodzą, że dołączanie zlogarytmowanych opóźnień zmiennych Google Trends poprawia siłę predykcyjną modeli dla Czech i Węgier, podczas gdy rezultaty dla Polski okazują się niekorzystne. Słowo kluczowe, którego używają autorzy w modelu dla Polski, to *praca*. Zwracają jednak uwagę, że wybór tego terminu mógł nie być najlepszy, ponieważ polscy pracownicy charakteryzują się wysokim stopniem mobilności międzynarodowej, co uwidacznia się w szczególnie częstym wyszukiwaniu w Polsce brytyjskich portali z ogłoszeniami o pracę.

Nowcasting stopy bezrobocia w Polsce z wykorzystaniem danych Google został podjęty również przez Anttonena (2018), który w badaniu wykorzystał model BVAR (*Bayesian vector autoregression*) z komponentem sezonowym i przeprowadził analizę dla poszczególnych krajów członkowskich UE. Nie jest jednak jasne, jakich terminów użył do wygenerowania szeregów czasowych wykorzystywanych w badaniu.

3. Metoda badania

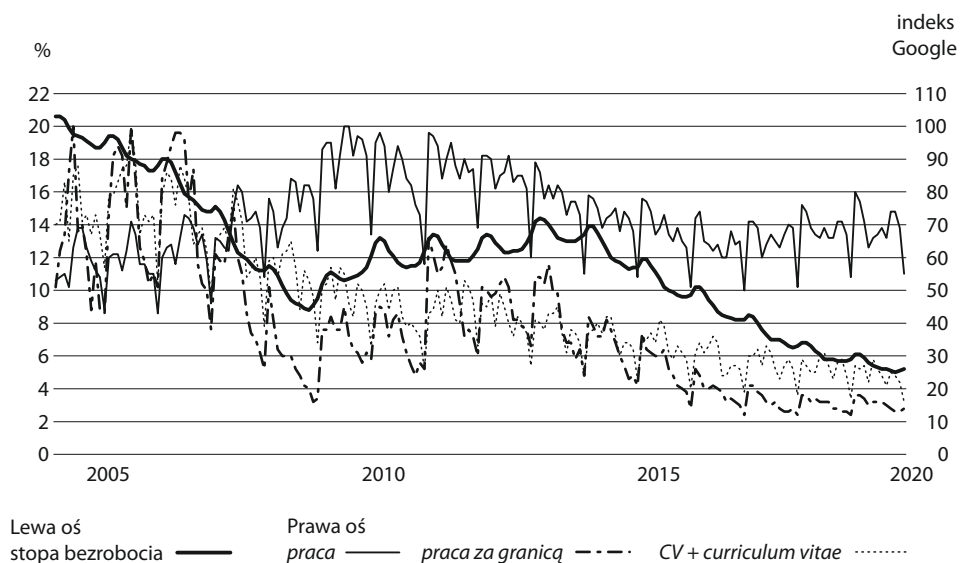
Główny Urząd Statystyczny regularnie publikuje oficjalne statystyki dotyczące stopy bezrobocia (definiowanej jako udział zarejestrowanych bezrobotnych w aktywnej zawodowo ludności cywilnej) według stanu na koniec miesiąca. Jednak statystyki te dostępne są z blisko miesięcznym opóźnieniem (dla przykładu statystyki dotyczące styczniowej stopy bezrobocia dostępne są pod koniec lutego). Innymi słowy, aktualna stopa bezrobocia nie jest znana.

Na wykresie przedstawiono dane dotyczące stopy bezrobocia rejestrowanego (GUS, 2021) w Polsce od stycznia 2004 r. do grudnia 2019 r. oraz dane Google Trends dla tego samego okresu.

Wysokość stopy bezrobocia rejestrowanego w Polsce podlegała w analizowanym okresie silnym fluktuacjom. W pierwszym kwartale 2004 r. wynosiła ponad 20%, ale wejście Polski do UE wiązało się z przyspieszeniem gospodarczym, a tym samym ze wzrostem popytu na pracę, co w konsekwencji prowadziło do szybkiego spadku bezrobocia (Bartosik, 2012). Wartość stopy bezrobocia spadała aż do października 2008 r., gdy osiągnęła poziom 8,8%. Od listopada 2008 r. następował ponowny, stopniowy wzrost bezrobocia, co wiązało się z globalnym kryzysem finansowym.

Stabilizacja sytuacji gospodarczej sprawiła, że od połowy 2014 r. bezrobocie zaczęło ponownie spadać. Rekordowo niski poziom stopy bezrobocia rejestrowanego odnotowano w październiku 2019 r., kiedy wyniosła ona 5%.

Wykres. Stopa bezrobocia rejestrowanego oraz wybrane indeksy Google Trends



Źródło: opracowanie własne z wykorzystaniem pakietu statystycznego R na podstawie: GUS (2021) i Google Trends.

Relatywnie szybki wzrost stopy bezrobocia pod koniec 2008 r. był trudny do przewidzenia. Nadzieję na uniknięcie generowania wyraźnie błędnych prognoz w przyszłości można pokładać w nowych źródłach danych, jak chociażby tych dotyczących wyszukiwania informacji w internecie.

Google Trends mierzy wolumen słów kluczowych bądź fraz, które są wyszukiwane przez użytkowników Google, a następnie porównuje ich popularność z innymi słowami kluczowymi lub frazami. Dokładne dane o samym wolumenie nie są publicznie dostępne. Zamiast tego Google Trends publikuje indeks, który wyliczany jest poprzez podzielenie liczby wyszukiwania danego terminu przez całkowitą liczbę zapytań użytkowników. Wyliczony ułamek skaluje się następnie do przedziału 0–100. Wartość 100 odpowiada maksymalnej liczbie wyszukiwań konkretnego terminu. Pozostałe wartości szeregu skalowane są zgodnie z wartością szczytową. Wynika to z prostego faktu: internet jest wciąż stosunkowo nową technologią i wzrostowy trend liczby użytkowników przekłada się na analogiczny wzrostowy trend w absolutnych wartościach większości zapytań.

Na potrzeby niniejszego badania wygenerowane zostały trzy szeregi czasowe Google Trends zobrazowane na wykresie, tj.: *praca*, *praca za granicą* oraz *CV + curriculum vitae* dla okresu od stycznia 2004 r. do grudnia 2019 r. Żadna z obserwacji dla wszystkich trzech indeksów nie wynosi 0, co oznacza, że liczba zapytań była zawsze większa od wartości granicznej. Wygenerowane indeksy Google, analogicznie do stopy bezrobocia rejestrowanego, są podane w częstotliwości miesięcznej oraz nie zostały uprzednio wyrównane sezonowo.

W celu określenia siły predykcyjnej, jaką niesie ze sobą włączenie do modelu zmiennych Google Trends, niezbędne jest w pierwszej kolejności skonstruowanie modelu, który może służyć jako benchmark. Ze względu na cel badania oraz sposób podejścia Pavlicka i Kristoufka (2015) w omawianym badaniu zbiór potencjalnych modeli, które mogłyby posłużyć jako punkt odniesienia, został ograniczony do klasy modeli autoregresyjnych (AR). Proste, jednowymiarowe modele autoregresyjne dość często wykorzystywane są jako benchmarki w literaturze z zakresu prognozowania szeregów czasowych.

W modelu autoregresyjnym (AR) zmienna zależna y_t zależy liniowo od swoich przeszłych wartości. Model autoregresyjny rzędu p można zapisać jako:

$$y_t = \mu + \phi_1 y_{t-1} + \phi_2 y_{t-2} + \dots + \phi_p y_{t-p} + \varepsilon_t, \quad (1)$$

gdzie:

y_t – wartość zmiennej prognozowanej w momencie t ,

y_{t-i} – opóźnione w czasie wartości zmiennej prognozowanej,

μ – stała modelu,

ϕ_j – parametry modelu,

$\varepsilon_t \sim N(0, \sigma)$.

Przed przystąpieniem do określenia odpowiedniego rzędu modelu AR niezbędne jest sprawdzenie założenia o stacjonarności zmiennych. W celu zbadania stacjonarności przeprowadzone zostały: rozszerzony test Dickeya-Fullera (ADF) oraz test Kwiatkowskiego-Phillipsa-Schmidta-Shina (KPSS). Testy te mają odwrotne hipotezy zerowe, w związku z czym tworzą komplementarną parę, powszechnie wykorzystywaną do testowania stacjonarności.

Statystyki testowe obu testów wyliczone zostały dla sześciu opóźnień (analogiczne rezultaty otrzymano dla mniejszych oraz większych rzędów – maksymalny rozważany rząd wynosił 13). Wyniki rozszerzonego testu ADF wskazywały na konieczność wykorzystania pierwszych różnic dla szeregu czasowego stopy bezrobocia rejestrowanego oraz szeregu *praca*. Wyniki testu KPSS sugerowały z kolei potrzebę zastosowania różnicowania wobec wszystkich zmiennych wykorzystywanych w badaniu.

Ze względu na wyniki obu testów do procesu modelowania wykorzystano pierwsze różnice wszystkich zmiennych. Różnicowanie stopy bezrobocia wyrażonej w procentach oraz indeksu Google Trends nie pozwala na prostą interpretację parametrów estymowanych modeli, jednak – jak wspomniano we wstępie – nie jest to celem tej pracy. Warto nadmienić, że interpretacja indeksu Google Trends jest trudna, nawet jeśli nie zostaje on poddany jakiegokolwiek transformacji, co wynika ze złożoności oraz niejasności metodyki jego wyliczania.

Do określenia rzędu opóźnień w modelu $AR(p)$ wykorzystuje się m.in. funkcję autokorelacji cząstkowej (*partial autocorrelation function* – PACF). Z wykresu wygenerowanego na potrzeby badania wynika, że pierwsze opóźnienia zróżnicowanej stopy bezrobocia miały relatywnie najwyższą wartość autokorelacji cząstkowej w porównaniu z pozostałymi opóźnieniami. Autokorelacja cząstkowa stała się statystycznie nieistotna na poziomie 5% dopiero dla opóźnień dalszych niż 13. Wykorzystanie modelu $AR(13)$ jako benchmarku nie wydaje się jednak zasadne, jeśli chodzi o zastosowanie go w nowcastingu. Włączenie parametru dla stałej wymagałoby oszacowania łącznie aż 14 parametrów, podczas gdy szereg stopy bezrobocia rejestrowanego liczy jedynie 192 obserwacje, co mogłoby prowadzić do nadmiernego dopasowania modelu do danych.

Co więcej, zbyt duża liczba parametrów byłaby konieczna do oszacowania również wówczas, gdy podstawy wyboru rzędu modelu $AR(p)$ nie stanowiłaby wizualna analiza wykresu funkcji autokorelacji cząstkowej, lecz byłoby nią kryterium informacyjne Akaike (*Akaike information criterion* – AIC), które jest jednym z formalnych sposobów wyboru pomiędzy modelami ekonometrycznymi o różnej liczbie predyktorów. Próba automatycznego dopasowania modelu autoregresyjnego do zróżnicowanego szeregu stopy bezrobocia rejestrowanego skutkowała bowiem wskazaniem modelu $AR(17)$ jako optymalnego. W mocy pozostaje zatem argument dotyczący zbyt małej liczby obserwacji w próbie.

W związku z powyższymi problemami z wyborem rzędu modelu $AR(p)$ na podstawie wykresu funkcji autokorelacji cząstkowej oraz AIC, wybór liczby opóźnień podstawowego modelu miał charakter arbitralny. Jako benchmark wykorzystano model $AR(1)$, którego użyli np. Choi i Varian (2009) w pionierskim artykule dotyczącym nowcastingu z wykorzystaniem Google Trends. Warto podkreślić, że autorzy korzystali w modelu ze zlogarytmowanego komponentu autoregresyjnego.

Montgomery i in. (1998) sugerują wybór modelu autoregresyjnego z wyłącznie jednym opóźnieniem do generowania krótkookresowych prognoz stopy bezrobocia. Prosty, tj. zawierający małą liczbę zmiennych objaśniających, model podstawowy wydaje się wystarczający do sprawdzenia, czy dane dotyczące wyszukiwanych fraz są przydatne w procesie prognozowania. Jeśli szeregi Google Trends nie okazałyby się przydatnym rozszerzeniem modelu $AR(1)$ w predykcji stopy bezrobocia, wówczas

mało prawdopodobne byłoby to, że stałyby się one cennym dodatkiem w bardziej skomplikowanych modelach ekonometrycznych. Co więcej, dane dotyczące wielkości bezrobocia w próbie z lat 2004–2019 charakteryzują się – jak wspomniano – dość dużą zmiennością. Nie jest do końca jasne, jak w świetle tej zmienności powinna być modelowana dynamika tego szeregu czasowego, ponieważ historyczne epizody mogłyby zbyt silnie oddziaływać na wartość zmiennej zależnej i tym samym zaburzać cały proces prognozowania.

Model autoregresyjny z opóźnieniem pierwszego rzędu został następnie rozszerzony o komponent sezonowy y_{t-12} . Analogiczne podejście stosują np. Choi i Varian (2009) oraz Tuhkuri (2016). Dodatkowo Montgomery i in. (1998) w badaniu stopy bezrobocia w Stanach Zjednoczonych zauważają, że charakteryzuje się ona sezonowością w długim okresie. Oprócz oparcia w literaturze dołączenie komponentu sezonowego sugerowała również wizualna analiza funkcji autokorelacji (*autocorrelation function* – ACF), która jest pomocna w identyfikacji dwóch składowych szeregów czasowych, tj. trendu oraz sezonowości. Wykres funkcji autokorelacji został wygenerowany dla zróżnicowanego szeregu stopy bezrobocia rejestrowanego, w związku z czym występowanie trendu nie było na nim widoczne, pomimo że szereg ten przed transformacją charakteryzował się trendem malejącym (o czym wspomniano przy opisie danych – bezrobocie w Polsce zaczęło szybko spadać w okresie poakcesyjnym). Wyraźne pozostały jednak wahania periodyczne o częstotliwości rocznej.

W tym miejscu należy zauważyć, że model AR(1) z komponentem sezonowym nie jest identyczny z modelem opisującym zmienność stopy bezrobocia w Polsce. W celu sprawdzenia, czy wybrany benchmark w zadowalającym stopniu opisuje proces generujący dane, przeprowadzono testy diagnostyczne. W pierwszej kolejności zweryfikowano założenie o stałości wariancji reszt, do czego wykorzystano test Breusch-Pagana (B-P). Hipoteza zerowa w tym teście mówi o homoskedastyczności składnika losowego analizowanego modelu, a hipoteza alternatywna – o jego heteroskedastyczności. Dla modelu AR(1) z komponentem sezonowym hipoteza zerowa została odrzucona przy wartości $p < 0,01$. W modelu służącym jako benchmark w tym badaniu mieliśmy do czynienia z tym samym problemem heteroskedastyczności. Kolejnym krokiem statystycznej diagnostyki modelu było przeprowadzenie testu Ljung-Boxa (L-B). Hipoteza zerowa tego testu mówi o braku autokorelacji składnika losowego, a hipoteza alternatywna – o występowaniu autokorelacji. Dla podstawowego modelu hipoteza zerowa testu L-B została odrzucona przy wartości $p < 0,01$ dla rzędu opóźnień równego 3, 6, 12 oraz 24. A zatem komponent sezonowy modelu nie był w stanie uwzględnić całej sezonowości występującej w zróżnicowanym szeregu czasowym stopy bezrobocia rejestrowanego. Nie jest to szczególnie zaskakujące, jeśli wziąć po uwagę wykres ACF oraz rząd opóźnień wskazywany przez AIC.

Model AR(1) z komponentem sezonowym estymowany był w niniejszym badaniu przy wykorzystaniu metody najmniejszych kwadratów. Jednak w kontekście wyników testów B-P oraz L-B nieodzowne wydawało się skorzystanie z odpornych błędów standardowych. Niezbędną korektę wyników estymacji przeprowadzono z wykorzystaniem błędów standardowych HAC (*heteroskedasticity and autocorrelation consistent*) zaproponowanych przez Newey i Westa (1987, 1994). Pozwalają one na uniknięcie konieczności doprecyzowania natury zautokorelowanego błędu, co jest wymagane podczas korzystania z alternatywnych estymatorów o mniejszej wariancji. Warto jeszcze raz podkreślić, że restrykcje nałożone na model podstawowy były znaczne. Aby sprawdzić wrażliwość wyników jego estymacji na zmianę liczby opóźnień, dodatkowo oszacowano modele AR(2)–AR(6) z niezmienionym komponentem sezonowym. Oceniono statystyczną istotność poszczególnych zmiennych, a same modele zostały porównane na podstawie AIC oraz Bayesowskiego kryterium informacyjnego Schwarza (*Bayesian information criterion* – BIC).

Przed przystąpieniem do rozszerzenia modelu podstawowego o zmienne Google Trends przetestowano przyczynowość w sensie Grangera (1969). Sprawdzenie, czy indeksy Google mogą być przydatne do monitorowania bieżącego stanu gospodarki, za pomocą tekstu Grangera zastosował np. Suhoj (2009). Sam test polega na weryfikacji, czy opóźnione wartości jednej zmiennej w modelu autoregresji wektorowej (*vector autoregression model* – VAR) są pomocne w prognozowaniu wartości innej zmiennej.

Przed przystąpieniem do rozszerzenia modelu AR(1) z komponentem sezonowym o egzogeniczne zmienne Google Trends warto podkreślić, że są one dostępne w czasie rzeczywistym, tzn. indeksy wyszukiwania, np. ofert pracy w styczniu, są dostępne w tym samym miesiącu. Jak wspomniano, oficjalne statystyki dotyczące stopy bezrobocia rejestrowanego w Polsce są publikowane przez GUS z prawie jednomiesięcznym opóźnieniem, co sprawia, że indeksy Google mają swoistą przewagę, jeśli chodzi o ich bieżącą dostępność. Dostępność Google Trends w momencie t , podczas gdy dane dotyczące bezrobocia dostępne są dopiero w momencie $t + 1$, skłania do użycia ich do monitorowania aktualnego poziomu agregatów. Wykorzystane w badaniu model benchmarkowy oraz modele rozszerzone o poszczególne indeksy Google przedstawiają się następująco:

$$\text{model 0: } \Delta y_t = \beta_{00} + \beta_{10}\Delta y_{t-1} + \beta_{20}\Delta y_{t-12} + e_t, \quad (2)$$

$$\text{model 1: } \Delta y_t = \beta_{01} + \beta_{11}\Delta y_{t-1} + \beta_{21}\Delta y_{t-12} + \beta_{31}\Delta GT_{1,t} + e_t, \quad (3)$$

$$\text{model 2: } \Delta y_t = \beta_{02} + \beta_{12}\Delta y_{t-1} + \beta_{22}\Delta y_{t-12} + \beta_{32}\Delta GT_{2,t} + e_t, \quad (4)$$

$$\text{model 3: } \Delta y_t = \beta_{03} + \beta_{13}\Delta y_{t-1} + \beta_{23}\Delta y_{t-12} + \beta_{33}\Delta GT_{3,t} + e_t, \quad (5)$$

gdzie:

Δy_t – pierwsze różnice wartości zmiennej prognozowanej w momencie t ,

Δy_{t-i} – opóźnione w czasie wartości zmiennej prognozowanej,

$\Delta GT_{j,t}$ – pierwsze różnice indeksów Google Trends w momencie t ,

$\beta_{k,m}$ – parametry modelu,

ΔGT – zróżnicowany indeks Google Trends,

e_t – składnik losowy modelu.

Modele rozszerzone o zmienną egzogeniczną GT były estymowane analogicznie do modelu podstawowego. W celu oszacowania parametrów zastosowano metodę najmniejszych kwadratów, a niezbędną korektę wynikającą z występowania problemów związanych z heteroskedastycznością oraz autokorelacją wykonywano przy użyciu błędów standardowych HAC, zaproponowanych przez Neweya i Westa (1987, 1994).

Wszystkie cztery modele w pierwszej kolejności były estymowane na całej dostępnej do badania próbie, tj. 192 obserwacjach. Na podstawie oszacowanych w ten sposób parametrów sprawdzono jakość dopasowania modeli do danych.

Na kolejnym etapie badania zbiór danych został podzielony na część treningową oraz testową. Podobne podejście zastosowali Pavlicek i Kristoufek (2015), którzy dane z okresu od stycznia 2004 r. do grudnia 2013 r. podzielili na zbiór treningowy, liczący 96 obserwacji (tj. od stycznia 2004 r. do grudnia 2011 r.), oraz zbiór testowy, liczący 24 obserwacje (tj. od stycznia 2012 r. do grudnia 2013 r.). Ze względu na to, że długość szeregów czasowych w badaniu przedstawianym w niniejszym artykule jest większa, zdecydowano się na wykorzystanie do nowcastingu nie dwóch, lecz trzech ostatnich lat. Tym samym w celu dopasowania modelu do danych sięgnięto po 156 obserwacji (tj. od stycznia 2004 r. do grudnia 2016 r.), a predykcje wygenerowano dla 36 obserwacji (tj. od stycznia 2017 r. do grudnia 2019 r.). Zastosowanie tego podziału pozwoliło pośrednio odpowiedzieć na pytanie, które z analizowanych szeregów Google Trends mogą być przydatne w monitorowaniu bieżącego poziomu stopy bezrobocia rejestrowanego, a także ocenić wielkość oraz statystyczną istotność tej poprawy. W związku z tym nie przeprowadzano dodatkowej oceny jakości modeli oszacowanych na danych ze zbioru treningowego. Nacisk na tym etapie został położony na samo prognozowanie z wykorzystaniem wciąż relatywnie nowego źródła danych.

Trafność prognoz wygenerowanych na zbiorze testowym sprawdzono za pomocą dwóch miar błędów *ex post*, tj. średniej wartości bezwzględnej błędu (*mean absolute error* – MAE) oraz pierwiastka średniego kwadratu błędu (*root mean square error* – RMSE). Z tych samych miar odchylenia rzeczywistych wartości realizacji zmiennej

zależnej od obliczonych prognoz korzysta, w kontekście prognozowania z wykorzystaniem Google Trends, np. Önder (2017).

Porównanie dokładności, przy użyciu miar błędów *ex post*, prognoz *out-of-sample* modelu podstawowego oraz modeli rozszerzonych o zmienne Google Trends pozwala na ocenę, które z szeregów Google są przydatne w procesie monitorowania bieżącego poziomu stopy bezrobocia. Mówiąc precyzyjniej, jeśli miary wyliczone dla modeli rozszerzonych przyjmują wartości mniejsze niż miary wyliczone dla modelu AR(1) z komponentem sezonowym, to można stwierdzić, czy indeksy Google są przydatne w procesie nowcastingu oraz jak znacząca jest poprawa predykcji związana z włączeniem ich do modelu. Niemniej konieczne wydaje się sprawdzenie, czy różnica w dokładności wygenerowanych prognoz jest istotna statystycznie. W związku z tym w badaniu wykorzystano test Diebolda-Mariano (D-M) zaproponowany przez Diebolda i Mariano (2002) oraz Westa (1996). Test D-M pozwala na porównanie dwóch lub więcej alternatywnych prognoz. Hipoteza zerowa testu mówi, że różnica w dokładności predykcji jest nieistotna (Diebold i Mariano, 2002).

4. Wyniki badania

Jak stwierdzono w części dotyczącej metody badania, model podstawowy AR(1) z komponentem sezonowym jest dość restrykcyjny. Niemniej w porównaniu z modelami AR(2)–AR(6) z analogicznymi komponentami sezonowymi Δy_{t-12} jego własności statystyczne wydają się korzystne. Obliczone kryteria informacyjne AIC i BIC sugerowały wybór modelu autoregresyjnego z sześcioma opóźnieniami, chociaż dla AR(6) istotne statystycznie na poziomie 5% okazały się jedynie zmienne Δy_{t-1} oraz Δy_{t-12} , czyli składowe wykorzystanego modelu podstawowego. Testy L-B oraz B-P przeprowadzone dla AR(6) z komponentem sezonowym również nie wykazały, że mógłby on zaradzić problemom związanym z autokorelacją oraz heteroskedastycznością składnika losowego. Modele AR(2)–AR(5) cechowały się wyższymi wartościami AIC niż model podstawowy. W związku z tym – oraz z uwagi na zbyt dużą liczbę zmiennych niezbędnych do oszacowania w modelach wskazywanych przez wykres PACF oraz na automatyczny dobór modelu na podstawie AIC – za punkt odniesienia posłużył model AR(1) z komponentem sezonowym. Warto w tym miejscu nadmienić, że ten model charakteryzował się wysoką wartością współczynnika determinacji $R^2 = 0,78$. Goel i in. (2010) wskazują, że w wielu przypadkach prosty model autoregresyjny z jednym opóźnieniem oraz z komponentem sezonowym może wyjaśniać ponad 90% zmienności zmiennej zależnej. Dlatego z dużą dozą ostrożności należy podchodzić do analizy nowych źródeł danych opierającej się jedynie na jakości dopasowania modeli do danych. Większe znaczenie ma trafność predykcji generowanych przez alternatywne modele.

Hipoteza mówiąca, że szeregi czasowe Google Trends mogą zostać użyte jako rozszerzenie modelu podstawowego do prognozowania stopy bezrobocia rejestrowanego w Polsce, została zweryfikowana pozytywnie. W teście Grangera bowiem hipoteza zerowa mówiąca, że zróżnicowany indeks Google nie jest przyczyną w sensie Grangera zróżnicowanej stopy bezrobocia, została odrzucona dla wszystkich trzech szeregów (*praca*, *praca za granicą* oraz *CV + curriculum vitae*) na poziomie istotności wynoszącym 1%. Niemniej relacja kauzalna (w rozumieniu probabilistycznej koncepcji przyczynowości) nie jest w tym przypadku jednoznaczna. Hipoteza zerowa mówiąca, że zróżnicowana stopa bezrobocia nie jest przyczyną w sensie Grangera zróżnicowanego indeksu Google, również została odrzucona dla wszystkich trzech szeregów Google Trends na poziomie istotności wynoszącym 5%. Tym samym do wyników testu należy podchodzić ostrożnie.

Następnie model podstawowy AR(1) z komponentem sezonowym został rozszerzony o zmienne egzogeniczne Google. W tablicy przedstawiono oszacowania parametrów poszczególnych modeli oraz podstawowe miary ich dopasowania do danych. Jak wspomniano, wstępne dopasowanie oraz ocena statystycznych charakterystyk wszystkich czterech modeli zostały przeprowadzone dla wartości wyliczonych na podstawie wszystkich 192 obserwacji z próby (tj. dla danych od stycznia 2004 r. do grudnia 2019 r.), a samą estymację przeprowadzono z wykorzystaniem metody najmniejszych kwadratów oraz skorygowano przy użyciu błędów standardowych HAC.

Tablica. Oszacowania parametrów modeli oraz miary ich dopasowania do danych

Zmienne objaśniające	Modele			
	0	1	2	3
<i>const.</i>	-0,0018 (0,0244)	-0,0026 (0,0254)	-0,0034 (0,0237)	-0,0018 (0,0262)
Δy_{t-1}	0,2899** (0,1070)	0,2955** (0,1063)	0,2831** (0,1049)	0,2896** (0,1044)
Δy_{t-12}	0,6653*** (0,0800)	0,6491*** (0,0862)	0,6385*** (0,0868)	0,6650*** (0,0808)
ΔGT_t	.	0,0019 (0,0012)	0,0034* (0,0016)	0,0001 (0,0006)
Skorygowany R^2	0,7783	0,7806	0,7844	0,7770
<i>AIC</i>	-174,0290	-174,8926	-178,0551	-172,0383
<i>BIC</i>	-161,2795	-158,9556	-162,1181	-156,1013

Uwaga. W nawiasach podano błędy standardowe. ***, **, * – zmienne istotne odpowiednio na poziomie: 1%, 5% i 10%. R^2 – współczynnik determinacji. Modele: 0 – podstawowy, 1–3 – rozszerzone o zmienne egzogeniczne odpowiednio: *praca*, *praca za granicą*, *CV + curriculum vitae*.

Źródło: opracowanie własne z wykorzystaniem pakietu statystycznego R na podstawie danych GUS (2021).

Jak można zauważyć na podstawie danych zamieszczonych w tablicy, oszacowania parametrów przy zmiennych Google Trends okazują się większe od 0, co oznacza, że

wyszukiwane słowa kluczowe lub frazy są pozytywnie powiązane ze stopą bezrobocia rejestrowanego. Można więc stwierdzić, że poczynione we wstępie założenie, że pogarszająca się bądź zła sytuacja na rynku pracy jest główną przyczyną zwiększonej aktywności osób poszukujących zatrudnienia w internecie, okazało się trafne. Niemniej do tego stwierdzenia należy podchodzić z dużą rezerwą ze względu na to, że jedynie zmienna *praca za granicą* okazała się istotna statystycznie. Dla zmiennych *praca* oraz *CV + curriculum vitae* nie została odrzucona hipoteza zerowa testu *t*-Studenta mówiąca, że oszacowany parametr wynosi 0. Dodatkowo w modelach pozostawiono nieistotną statystycznie stałą, ponieważ jej usunięcie zasadniczo nie powodowało istotnych zmian w oszacowaniach parametrów oraz stosunków pomiędzy wyliczonymi miarami dopasowania modeli do danych. Należy jednak zauważyć, że w każdym z modeli silnie istotne statystycznie są składowe modelu AR(1) z komponentem sezonowym. Zmienna Δy_{t-12} okazała się dla wszystkich estymowanych modeli istotna statystycznie na poziomie istotności 1%, co potwierdza zasadność jej użycia. Warto również podkreślić, że wszystkie modele cechowały się wysoką wartością skorygowanego współczynnika determinacji R^2 , jednak dołączenie do modelu podstawowego szeregu *praca*, *praca za granicą* czy *CV + curriculum vitae* nieznacznie wpłynęło na jego wartość (przy czym dołączenie szeregu *CV + curriculum vitae* – negatywnie). Fakt ten można uznać za potwierdzenie tezy Goela i in. (2010), zgodnie z którą prosty model autoregresyjny z komponentem sezonowym wyjaśniał znaczną część zmienności zróżnicowanej stopy bezrobocia rejestrowanego. Warto też dodać, że wskazywały na to wyliczone wartości AIC oraz BIC, które jedynie nieznacznie różniły się pomiędzy poszczególnymi modelami 0–3, przy czym najmniejszymi wartościami obu kryteriów cechował się model ze zmienną egzogeniczną *praca za granicą*. Model ten okazał się tym samym najlepiej dopasowany do danych, a szereg Google Trends, który został do niego dołączony – jedynym istotnym statystycznie spośród trzech wykorzystanych w badaniu.

Na kolejnym etapie sprawdzono, które z modeli (dopasowanych tym razem na mniejszej liczbie obserwacji – zbiorze treningowym) pozwoliły na wygenerowanie najdokładniejszych prognoz zróżnicowanej stopy bezrobocia rejestrowanego w Polsce dla danych z okresu od stycznia 2017 r. do grudnia 2019 r. (tj. danych ze zbioru testowego). Nieznacznie mniejszymi, w porównaniu do modelu podstawowego, błędami prognoz *ex post* charakteryzował się jedynie model rozszerzony o wartości indeksu Google *praca za granicą*, zarówno jeśli chodzi o miarę MAE, jak i RMSE. Wartość MAE dla modelu AR(1) z komponentem sezonowym wynosiła 0,08, a wartość RMSE – 0,09. Z kolei dla modelu 2 – odpowiednio 0,07 oraz 0,09. Dołączenie szeregów *praca* oraz *CV + curriculum vitae* nie tylko nie poprawiło zatem miar dopasowania modelu do danych, lecz także nie przyczyniło się do zwiększenia trafności prognoz generowanych przy jego użyciu.

W teście D-M różnice w dokładności prognoz generowanych przez benchmark oraz model rozszerzony o wartości szeregu *praca za granicą* okazały się istotne statystycznie. Hipoteza zerowa mówiąca o braku tych różnic została odrzucona na 10-procentowym poziomie istotności. Z kolei dołączenie do modelu podstawowego indeksów Google Trends *praca* oraz *CV + curriculum vitae* nie wpłynęło na trafność predykcji. Wyniki testu D-M nie wskazywały bowiem na występowanie różnic pomiędzy prognozami wygenerowanymi przez benchmark a prognozami wygenerowanymi przez modele 1 i 3. Innymi słowy, modele ze zmiennymi egzogenicznymi *praca* oraz *CV + curriculum vitae* nie pogorszyły istotnie predykcji generowanych przez model AR(1) z komponentem sezonowym.

5. Podsumowanie

Celem badania przedstawionego w artykule było sprawdzenie, czy indeksy Google Trends poprawiają trafność predykcji autoregresyjnego modelu stopy bezrobocia rejestrowanego w Polsce. Zweryfikowano, czy szereg czasowy dotyczący wyszukiwania hasła *praca*, a także szeregi dotyczące wyszukiwania haseł *praca za granicą* oraz *CV + curriculum vitae* (niewykorzystywane wcześniej w literaturze) są użyteczne w procesie monitorowania poziomu tejże zmiennej makroekonomicznej.

Wyniki empiryczne pokazują, że słowo kluczowe *praca* jest mało przydatnym rozszerzeniem modeli wykorzystywanych do prognozowania stopy bezrobocia rejestrowanego w Polsce. Podobnie szereg Google Trends *CV + curriculum vitae* nie poprawia trafności predykcji generowanych przez model autoregresyjny.

Przydatny do monitorowania poziomu stopy bezrobocia jest natomiast indeks *praca za granicą*, co może wynikać z charakterystyk związanych z międzynarodową mobilnością siły roboczej w Polsce. Warto jednak zaznaczyć, że poprawa trafności okazała się niewielka. Niezbędne wydają się dalsze analizy doboru zmiennych egzogenicznych dotyczących aktywności użytkowników internetu odnoszącej się do poszukiwania zatrudnienia.

Niniejszy artykuł przedstawia zarówno ograniczenia, jak i potencjał związany z wykorzystywaniem wciąż relatywnie nowego źródła danych, jakim jest serwis Google Trends. Replikacja sposobów doboru indeksów Google, które są stosowane w literaturze anglojęzycznej z zakresu nowcastingu (a nawet szerzej – prognozowania), nie okazała się poprawnym podejściem do przeprowadzania analiz makroekonomicznych dla Polski. Jednakże uwzględnienie specyfiki języka polskiego oraz istotnych charakterystyk polskiej gospodarki może pomóc o wiele dokładniej przewidywać krajową sytuację ekonomiczną, a dzięki temu narzędzia polityki gospodarczej będą mogły być stosowane w bardziej adekwatny i skuteczny sposób.

Bibliografia

- Anttonen, J. (2018). *Nowcasting the Unemployment Rate in the EU with Seasonal BVAR and Google Search Data* (ETLA Working Papers No. 62). <http://pub.etla.fi/ETLA-Working-Papers-62.pdf>.
- Askatas, N., Zimmermann, K. F. (2009). Google Econometrics and Unemployment Forecasting. *Applied Economics Quarterly*, 55(2), 107–120. <https://doi.org/10.3790/aeq.55.2.107>.
- Bartosik, K. (2012). Popytowe i podażowe uwarunkowania polskiego bezrobocia. *Gospodarka Narodowa*, 260(11–12), 25–57. <https://doi.org/10.33119/GN/101003>.
- Bello-Orgaz, G., Jung, J. J., Camacho, D. (2016). Social big data: Recent achievements and new challenges. *Information Fusion*, 28, 45–59. <https://doi.org/10.1016/j.inffus.2015.08.005>.
- Blazquez, D., Domenech, J. (2018). Big Data sources and methods for social and economic analyses. *Technological Forecasting & Social Change*, 130, 99–113. <https://doi.org/10.1016/j.techfore.2017.07.027>.
- Buono, D., Mazzi, G. L., Kapetanios, G., Marcellino, M., Pappailias, F. (2017). Big data types for macroeconomic nowcasting. *EURONA – Eurostat Review on National Accounts and Macroeconomic Indicators*, (1), 93–145. <https://ec.europa.eu/eurostat/cros/system/files/euronaissue1-2017-art4.pdf>.
- Choi, H., Varian, H. R. (2009). *Predicting the Present with Google Trends*. https://www.google.com/googlegloss/pdfs/google_predicting_the_present.pdf.
- Cox, M., Ellsworth, D. (1997). Managing Big Data for Scientific Visualization. *ACM Siggraph*, 97, 21–38.
- D’Amuri, F., Marcucci, J. (2012). *The predictive power of Google searches in forecasting unemployment* (Bank of Italy Working Papers No. 891). https://www.bancaditalia.it/pubblicazioni/temi-discussione/2012/2012-0891/en_tema_891.pdf?language_id=1.
- Diebold, F. X., Mariano, R. S. (2002). Comparing Predictive Accuracy. *Journal of Business & Economic Statistics*, 20(1), 134–144. <https://doi.org/10.1198/073500102753410444>.
- Ettredge, M., Gerdes, J., Karuga, G. (2005). Using Web-based Search Data to Predict Macroeconomic Statistics. *Communications of the ACM*, 48(11), 87–92. <https://doi.org/10.1145/1096000.1096010>.
- Giannone, D., Reichlin, L., Small, D. (2008). Nowcasting: The real-time informational content of macroeconomic data. *Journal of Monetary Economics*, 55(4), 665–676. <https://doi.org/10.1016/j.jmoneco.2008.05.010>.
- Goel, S., Hofman, J. M., Lahaie, S., Pennock, D. M., Watts, D. J. (2010). Predicting consumer behavior with Web search. *Proceedings of the National Academy of Sciences of the United States of America*, 107(41), 17486–17490. <https://doi.org/10.1073/pnas.1005962107>.
- Główny Urząd Statystyczny. (2018). *Informacja o rozmiarach i kierunkach czasowej emigracji z Polski w latach 2004–2017*. https://stat.gov.pl/download/gfx/portalinformacyjny/pl/defaultaktualnosci/5471/2/11/1/informacja_o_rozmiarach_i_kierunkach_czasowej_emigracji_z_polski_2004-2017.pdf.
- Główny Urząd Statystyczny. (2021). *Stopa bezrobocia rejestrowanego w latach 1990–2021*. <https://stat.gov.pl/obszary-tematyczne/rynek-pracy/bezrobocie-rejestrowane/stopa-bezrobocia-rejestrowanego-w-latach-1990-2021,4,1.html?pdf=1>.
- Granger, C. W. J. (1969). Investigating causal relations by econometric models and cross-spectral methods. *Econometrica*, 37(3), 424–438. <https://doi.org/10.2307/1912791>.

- Kapetanios, G., Papailias, F. (2018). *Big Data & Macroeconomic Nowcasting: Methodological Review* (ESCoE Discussion Paper 2018-12). <https://escoe-website.s3.amazonaws.com/wp-content/uploads/2020/07/13161005/ESCoE-DP-2018-12.pdf>.
- Laney, D. (2001). 3D Data Management: Controlling Data Volume, Velocity, and Variety. Application Delivery Strategies. <http://blogs.gartner.com/doug-laney/files/2012/01/ad949-3D-Data-Management-Controlling-Data-Volume-Velocity-andVariety.pdf>.
- Montgomery, A. L., Zarnowitz, V., Tsay, R. S., Tiao, G. C. (1998). Forecasting the U.S. Unemployment Rate. *Journal of American Statistical Association*, 93(442), 478–493. <https://doi.org/10.1080/01621459.1998.10473696>.
- Newey, W. K., West, K. D. (1987). A Simple, Positive Semi-Definite, Heteroskedasticity and Autocorrelation Consistent Covariance Matrix. *Econometrica*, 55(3), 703–708. <https://doi.org/10.2307/1913610>.
- Newey, W. K., West, K. D. (1994). Automatic Lag Selection in Covariance Matrix Estimation. *Review of Economic Studies*, 61(4), 631–653. <https://doi.org/10.2307/2297912>.
- Önder, I. (2017). Forecasting Tourism Demand with Google Trends: Accuracy Comparison of Countries vs. Cities. *International Journal of Tourism Research*, 19(6), 648–660. <https://doi.org/10.1002/jtr.2137>.
- Pavlicek, J., Kristoufek, L. (2015). Nowcasting Unemployment Rates with Google Searches: Evidence from the Visegrad Group Countries. *PLoS ONE*, 10(5), 1–11. <https://doi.org/10.1371/journal.pone.0127084>.
- Suhoy, T. (2009). *Query Indices and a 2008 Downturn: Israeli Data* (Bank of Israel Discussion Paper No. 2009.06). <https://www.boi.org.il/deptdata/mehkar/papers/dp0906e.pdf>.
- Tuhkuri, J. (2016). *Forecasting Unemployment with Google Searches* (ETLA Working Papers No. 35). <https://www.econstor.eu/bitstream/10419/201250/1/ETLA-Working-Papers-35.pdf>.
- West, K. D. (1996). Asymptotic inference about predictive ability. *Econometrica*, 64(5), 1067–1084. <https://doi.org/10.2307/2171956>.

Wspomnienie o Władysławie Wiesławie Łagodzińskim

Remembering Władysław Wiesław Łagodziński

Droga zawodowa

Władysław Wiesław Łagodziński (1937–2020) rozpoczął pracę w Głównym Urzędzie Statystycznym w 1968 r., kilka lat po ukończeniu studiów bibliotekoznawczych na Uniwersytecie Warszawskim, gdzie studiował również socjologię. W pierwszym okresie pracy zawodowej był zatrudniony na stanowisku starszego radcy w Departamencie Statystyki Oświaty, Kultury i Spraw Socjalnych, a następnie pełnił funkcję wicedyrektora w Departamencie Badań Społecznych i Demograficznych. Początkowo zajmował się przede wszystkim szeroko rozumianą statystyką kultury. Publikował, głównie na łamach „Wiadomości Statystycznych”, artykuły poświęcone tej tematyce, co zostało docenione nadaniem mu honorowej odznaki resortowej Zasłużony Działacz Kultury. Drugim ważnym obszarem jego zainteresowań w latach 70. i 80. była statystyka turystyki.

Od 1990 r. aktywnie uczestniczył w badaniach warunków bytu ludności, w tym w badaniach gospodarstw domowych. Miał też udział we wdrażaniu Zintegrowanego Systemu Badań Gospodarstw Domowych w GUS i przyczynił się do rozwoju statystyki społecznej. Artykuły, które publikował w tym czasie na łamach „WS”, i opracowania jego autorstwa zamieszczane w publikacjach GUS dotyczyły bezrobocia, pomocy społecznej dla rodzin z dziećmi, życia dzieci w obszarach pogranicza, zróżnicowania warunków życia ludności oraz warunków życia i potrzeb młodzieży. Problematyka statystyki społecznej, zwłaszcza badanie ubóstwa i wykluczenia społecznego, zawsze była dla niego bardzo ważna.

W 1993 r. został rzecznikiem prasowym GUS i zastępcą dyrektora nowo utworzonego Departamentu Informacji. Funkcję rzecznika pełnił do przejścia na emeryturę w 2004 r. Promował przede wszystkim aktywną politykę informacyjną GUS i był znakomitym popularyzatorem statystyki publicznej w mediach. Rozumiał dziennikarzy i wszystkich poszukujących informacji statystycznej i starał się im pomóc w dotarciu do danych gromadzonych przez GUS. Stworzył system równorzędnego, równoczesnego i równego dostępu do informacji statystycznej w Polsce. Pod jego kierunkiem opracowano pierwsze wydane drukiem zasady udzielania informacji. Ponadto odegrał znaczącą rolę w promowaniu Narodowego Spisu Powszechnego 2002 i 2011 oraz Powszechnego Spisu Rolnego 2002 i 2010, a także historii spisów powszechnych na ziemiach polskich.

Jego zasługą było utworzenie w GUS w 1993 r. Centralnego Informatorium Statystycznego (CIS). Wcześniej informacji można było poszukiwać jedynie w Centralnej Bibliotece Statystycznej (CBS), gdzie funkcjonował mikroskopijny, dwuosobowy Dział Informacji Statystycznej. Od 1993 r. CIS współpracuje z Działem Informacji i Bibliografii CBS przy obsłudze zapytań użytkowników danych statystycznych (za udostępnianie danych współczesnych odpowiada CIS, a za dane historyczne – CBS).

Od 1993 r. Władysław Wiesław Łagodziński był członkiem Kolegium Redakcyjnego „Wiadomości Statystycznych”. Wyróżniał się zaangażowaniem w dyskusjach i chętnie służył swoją wiedzą socjologiczną i statystyczną w ocenie artykułów zgłaszanych do publikacji.



Fot. Archiwum US w Warszawie

1994 r., Łotwa, spotkanie zespołu pod kierownictwem wiceprezesa GUS Jana Kordosa (drugi z prawej) w sprawie kontraktu Banku Światowego na zaprojektowanie badań warunków życia gospodarstw domowych na Łotwie

W latach 1995–2004 pełnił kolejno funkcje: doradcy Prezesa, dyrektora Departamentu Analiz i Udostępniania Informacji, zastępcy dyrektora Departamentu Analiz i Udostępniania Informacji i zastępcy dyrektora Sekretariatu Prezesa. W tym okresie uczestniczył aktywnie m.in. w pracach komisji do spraw ochrony tajemnicy statystycznej, co miało zasadnicze znaczenie przy udostępnianiu danych opracowanych w systemie statystyki publicznej.

Od 1999 r. redagował dodatek do „Rzeczpospolitej”, w którym zamieszczał felietony dotyczące m.in. statystyki danej dziedziny życia, tajemnicy statystycznej, metodologii badań czy nowych przedsięwzięć badawczych. Łącznie opublikował ponad 130 felietonów.

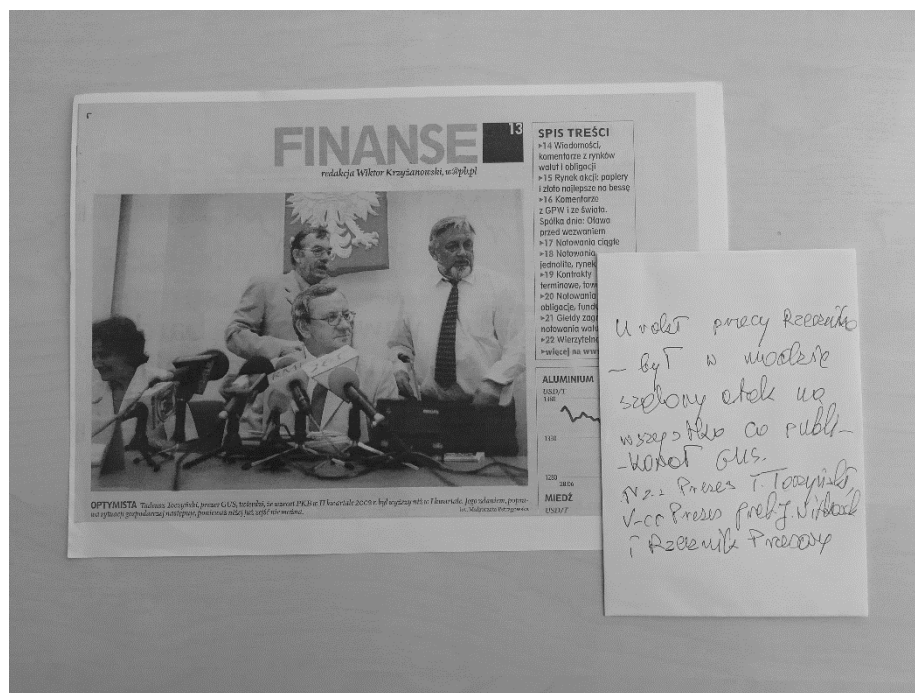


Fot. Archiwum Z. Kozłowskiej

2002 r., Jachranka, narada kadry kierowniczej GUS i urzędów statystycznych (notabene imieniny Wiesława)

W latach 2004–2006 pracował w Urzędzie Statystycznym (US) w Warszawie, gdzie prowadził prace nad rozwojem i popularyzacją statystyki oraz nad historią powojennej Warszawy w świetle danych liczbowych.

W latach 2006–2011 ponownie został rzecznikiem prasowym GUS, a jednocześnie – z upoważnienia prezesa GUS – nadzorował prace Departamentu Informacji, Zakładu Wydawnictw Statystycznych i CBS. Z jego inicjatywy organizowano wówczas tematyczne seminaria statystyczne, wystawy dokumentów historycznych poświęconych statystyce i opracowywano wydawnictwa jubileuszowe statystyki publicznej. Od 2009 r. był radcą generalnym Prezesa GUS. Formalnie zakończył pracę w GUS w 2011 r., ale nadal pozostawał czynnym statystykiem, a jego zainteresowania skupiały się na edukacji statystycznej i popularyzacji statystyki. Swój pokaźny księgozbiór przekazał CBS.



2002 r., konferencja prasowa prezesa GUS. Władysław Wiesław Łagodziński zachował ten wycinek prasowy z „Pulsu Biznesu” w swoim prywatnym archiwum, z komentarzem: „Uroki pracy rzecznika – był w modzie szalony atak na wszystko, co publikował GUS”.

Na zdjęciu również (od lewej): wiceprezes Halina Dmochowska, wiceprezes Janusz Witkowski, prezes Tadeusz Toczyński

W ostatnich latach życia zorganizował, przy współpracy pracowników US w Warszawie i CBS, kilkadziesiąt naukowych seminariów statystycznych, które odbywały się w gmachu US i były transmitowane do czytelni biblioteki oraz do urzędów statystycznych. Dzięki temu każdorazowo uczestniczyło w nich w czasie rzeczywistym kilkadziesiąt osób, a pozostali zainteresowani mogli odtworzyć nagrania udostępniane w internecie. Niestrudzenie udzielał wywiadów w prasie, radiu i telewizji, popularyzując statystykę publiczną i wyniki prac prowadzonych w jej ramach. Ponadto brał udział w pracach komisji egzaminacyjnych w zawodach szczebla okręgowego Olimpiady Statystycznej organizowanej przez GUS. W latach 2015–2017 aktywnie uczestniczył w pracach zespołu redakcyjnego serii *Historia Polski w liczbach*. Opracowano wówczas tomy 4 (*Statystyka Polski. Dawniej i dzisiaj*) i 5 (*Polska 1918–2018*).

Równoległe do pracy w GUS prowadził szeroką działalność w Polskim Towarzystwie Statystycznym (PTS). Po reaktywacji PTS 16 kwietnia 1981 r. brał aktywny udział w staraniach o uzyskanie przez Towarzystwo prawa do prowadzenia działal-

ności gospodarczej, które pozwoliłyby na zdobycie niezależności finansowej. W latach 2010–2018 pełnił funkcję wiceprezesa Rady Głównej PTS, a od 2011 r. – przewodniczącego Rady Oddziału Warszawskiego PTS. W 2019 r. został honorowym członkiem PTS.



Fot. Archiwum Z. Kozłowskiej

2008 r., Zamek Królewski w Warszawie, uroczystość z okazji 90-lecia GUS

Był inicjatorem powołania Biura Badań i Analiz Statystycznych (BBiAS), które utworzono 28 lutego 1986 r. na mocy uchwały Rady Głównej PTS, i pełnomocnikiem ds. organizacji biura, a następnie objął nad nim kierownictwo. Jego rola jako statystyka i kierownika była nie do przecenienia w rozwoju tej jednostki. W ramach BBiAS przeprowadzono wiele badań, m.in. Diagnozę Społeczną (osiem edycji: 2000, 2003, 2005, 2007, 2009, 2011, 2013 i 2015); badanie nastrojów społecznych dla Instytutu Studiów Społecznych Uniwersytetu Warszawskiego (w latach 1993–1994); badanie DS-44, czyli monitoring warunków życia ludności, dla Ministerstwa Pracy i Polityki Społecznej (1995); badania postaw społeczno-religijnych ludności Polski dla Instytutu Statystyki Kościoła Katolickiego (1997 i 2009); wspólnie z CBOS badanie *Warunki, problemy i strategie przetrwania*; badanie *Jacy jesteśmy?* (samookreśle-

nie warunków życia), zlecone przez Wydział Ekonomiczny Uniwersytetu Warszawskiego (2009), oraz badanie *Generation* (2010).



Fot. Archiwum US w Warszawie

2008 r., Ustrzyki Górne, spotkanie integracyjne z okazji 20-lecia współpracy polskich i słowackich statystyków. Na pierwszym planie delegacja z US w Preszowie, w głębi Marek Cierpień-Wolan, dyrektor US w Rzeszowie

Oprócz badań BBiAS prowadziło również działalność redakcyjną i wydawniczą, szkoleniową, konsultacyjną, seminaryjną i popularyzacyjną. W latach 1987–1992 zarejestrowano ponad 100 wykładów wybitnych polskich uczonych, m.in. profesorów Ryszarda Zasepy, Wiesława Sadowskiego, Jana Kordosa i Leszka Zienkowskiego. W 1990 r. Rada Główna PTS zainicjowała wydawanie serii Tłumaczenia, pod redakcją Władysława Wiesława Łagodzińskiego, w której publikowane były przekłady wydawnictw metodyczno-szkoleniowych w celu promocji najwybitniejszych osiągnięć statystyki światowej w środowisku polskich statystyków.

Władysław Wiesław Łagodziński z dużym zaangażowaniem na różnych frontach promował polską statystykę. Z okazji 200-lecia polskiej statystyki i 75-lecia GUS przy aktywnym udziale jego i innych członków PTS zorganizowano sześć konferencji: 29 czerwca 1993 r. w Warszawie, 22–23 września 1993 r. w Poznaniu, 24–25 maja 1993 r. w Mogilanach k. Krakowa, 25–26 października 1993 r. w Gdańsku, 27–29 września 1993 r. w Jachrance. Z ramienia PTS odpowiadał za przygotowanie sześciu

tomów zawierających referaty z tych konferencji, opublikowanych w serii Biblioteka Wiadomości Statystycznych (t. 42 – *Rozwój myśli i instytucji statystycznych na ziemiach polskich*, t. 43 – *Wyzwania polskiej statystyki*, t. 44 – *Rozwój metodologii badań statystycznych w Polsce*, t. 45 – *Statystyka regionalna*, t. 46 – *Obchody jubileuszowe 200-lecia statystyki polskiej – 75-lecia Głównego Urzędu Statystycznego*, t. 47 – *Podstawowe zasady statystyk oficjalnych oraz ich prawnych i etycznych aspektów w krajach w okresie przejściowym*).

Warto zwrócić uwagę na publikacje jego autorstwa dotyczące działalności PTS: *Dekalog zasad statystyk oficjalnych jako przesłanki do kształtowania się nowego ładu informacyjnego oraz Ewolucja badań warunków życia ludności i jej odniesienia do transformacji społeczno-ekonomicznej w Polsce*, które ukazały się w okolicznościowym wydawnictwie *Obchody jubileuszowe 200-lecia statystyki polskiej i 75-lecia Głównego Urzędu Statystycznego* (1995), rozdziały poświęcone Biuru Badań i Analiz Statystycznych przy Radzie Głównej PTS w pracach *Polskie Towarzystwo Statystyczne 1912–1992* (1992) oraz *Polskie Towarzystwo Statystyczne 1912–2012* (2012), a także artykuły opublikowane na łamach wydań specjalnych „VIP Magazyn”: *Finanse Polskiego Towarzystwa Statystycznego* (2014) i *Diagnoza Społeczna 2000–2013* (2014).

Władysław Wiesław Łagodziński został odznaczony Krzyżem Kawalerskim Orderu Odrodzenia Polski, Złotym Krzyżem Zasługi oraz wieloma odznaczeniami resortu statystyki.

Zmarł 18 grudnia 2020 r. w Warszawie w wieku 83 lat.

Dominik Rozkrut
Bożena Łazowska
Czesław Domański
Włodzimierz Okrasa

W oczach przyjaciół i współpracowników

Jak pisać w czasie przeszłym o Wieśku, który – jak uświadomiło mi jego odejście – był obecny w moim życiu od 30 lat? Zetknęłam się z nim na początku lat 90., kiedy już po kilku latach zajmowania się problematyką ludnościową czułam urok demografii i intelektualną obietnicę badań nad procesami ludnościowymi. Pierwsze spotkanie dotyczyło badań gospodarstw domowych prowadzonych przez GUS – chodziło o poszerzenie informacji o zdarzeniach demograficznych doświadczanych przez członków gospodarstw domowych. Odniosłam złe wrażenie z tej rozmowy; uznałam rozmówcę za osobę arogancką i obcesową. Najchętniej zrezygnowałabym z dalszych kontaktów, ale bardzo zależało mi na poszerzeniu zakresu danych. Podczas kolej-

nych spotkań, na których rozmawialiśmy nie tylko o badaniach gospodarstw domowych, zmieniałam zdanie o Wieśku. Przekonało mnie jego podejście do badań realizowanych przez GUS – a był to czas określania zadań statystyki publicznej w nowej rzeczywistości gospodarczej i społecznej. Jemu po prostu zależało, by szukać lepszych rozwiązań, i patrzył na to jako socjolog i statystyk. Był także miłośnikiem demografii, której się uczył u prof. Jerzego Z. Holzera. Ujęły mnie jego wielki szacunek i sympatia do Profesora. Odkryłam też poczucie humoru Wieśka, ironię, barwność języka, choć niejednokrotnie przeszkadzała mi dosadność jego krytyki, nierzadko prowokacyjnej. Był ciekawym rozmówcą.

Pod koniec lat 90. rozpoczęliśmy prace nad panelowym badaniem *Diagnoza Społeczna*, której był pomysłodawcą. Wspólna praca w tym projekcie i jego zrozumienie dla inicjatyw badawczych podejmowanych w Instytucie Statystyki i Demografii, zarówno krajowych, jak i międzynarodowych, sprawiły, że nasze relacje zmieniły się z koleżeńskich na przyjacielskie.

Wiesiek był pełnym ciekawości, wnikliwym obserwatorem procesów społecznych, który zarażał innych swoją energią i pasją do tworzenia danych statystycznych właściwych współczesnym przemianom. Jestem wdzięczna za jego otwartość na pomysły badawcze i sprzyjanie nowym badaniom prowadzonym przez demografów Instytutu Statystyki i Demografii po 2000 r. Był także nieustannie aktywny w przekazywaniu wiedzy o procesach demograficznych szerokiemu gronu odbiorców. Dzięki jego uporczywym naleganiom zespół pracowników instytutu i GUS uczył demografii w listopadzie 2019 r. w Szkole Zdrowia Publicznego Centrum Medycznego Kształcenia Podyplomowego. Wiesiek również należał do tego zespołu. Pamiętam długą rozmowę z nim po zakończeniu zajęć, pełną niepokoju o sytuację demograficzną Polski.

Był życzliwy młodym badaczom współpracującym przy kolejnych edycjach *Diagnozy Społecznej* w latach 2000–2015. Jego profesjonalizm, życzliwość i bezpośredniość z jednej strony, a poczucie humoru i ironia z drugiej zjednywały mu uznanie i szacunek oraz głęboką sympatię osób, które miały okazję się z nim zetknąć. Także podczas pełnienia przez wiele lat funkcji rzecznika prasowego GUS. Oddają to m.in. wpisy na profilu Patrycji Maciejewicz, dziennikarki „Gazety Wyborczej”, na Twitterze, będące reakcją na przekazaną przez nią wieść o śmierci Wieśka: „Totalnie wyjątkowy, znający instytucję i materię od podszewki, niemalostkowy, szczery, błyskotliwy i dusza towarzystwa”; „To był rzecznik instytucja, a nie rzecznik instytucji”; „Pamiętam Pana Wiesława, zawsze miał czas na rozmowę z dziennikarzem”; „Pamiętam, miał czas dla dziennikarza, ale też i zwykłego studenta zainteresowanego statystyką. RIP”. A funkcję rzecznika Wiesiek przestał pełnić w roku 2011...

Co mogę jeszcze powiedzieć o Wieśku z perspektywy środowiska demograficznego? Jak mówiłam podczas pogrzebu, wielu badaczy procesów ludnościowych zetknę-

ło się z nim na wczesnych etapach swojej biografii akademickiej, podczas pierwszych kontaktów dotyczących współpracy z GUS. Dla nas był to człowiek instytucja, nie tylko dlatego że angażował się w rozwój statystyki publicznej i jej promocję w Polsce, zwłaszcza od lat 90., kiedy pełnił odpowiedzialne funkcje w GUS i PTS. Słowo „zaangażowanie” nie oddaje istoty jego wieloletniej aktywności – najwłaściwszym określeniem wydaje się: „służył rozwojowi statystyki publicznej w Polsce”. I nie jest to określenie na wyrost.

O tym, że Wiesiek, jako socjolog współpracujący z Uniwersytetem Warszawskim, był aktywnym uczestnikiem protestów marcowych w 1968 r., dowiedziałam się dopiero po jego śmierci. Ja w tym czasie jako studentka Szkoły Głównej Planowania i Statystyki przechodziłam gwałtowny kurs dojrzewania obywatelskiego. Bardzo żałuję, że nie porozmawialiśmy o tym formatywnym doświadczeniu...

Jestem wdzięczna Wiesiowi za jego przykład życia człowieka niepokornego, działającego z pasją na rzecz dobra wspólnego, czego nie zmieniły trudne doświadczenia. Człowieka wrażliwego, ale też do bólu szczerego w krytyce.

Straciłam przyjaciela, który przypominał mi o powinnościach środowiska akademickiego. Staram się w miarę swoich sił je wypełniać, ale już nie mogę go zapytać, co o tym sądzi.

Irena E. Kotowska

*
* *

Wiesława poznałem ponad 30 lat temu. Dostałem pierwszy duży grant na badanie jakości życia Polaków w czasie transformacji systemowej. Wiedziałem, że jedyną profesjonalną siatką ankieterską jest ta GUS-owska. Umówiłem się na spotkanie z Wieskiem. On od razu zapalił się do projektu. Ze spotkania wyszedłem przekonany, że rozmawiałem raczej z prezesem GUS niż z jego rzecznikiem prasowym. Po pierwszym dwuletnim grantie był drugi dwuletni, a następnie trzyletni. Wiesiek zajmował się przez siedem lat organizacją badań w terenie. I tak dobrnęliśmy do końca wieku. W tym czasie byliśmy już przyjaciółmi, a nawet więcej – zaprzyjaźniły się nasze rodziny.

Gdy granty się skończyły, Wiesiek wpadł na genialny pomysł, żeby robić duże interdyscyplinarne badania panelowe dające panoramiczny obraz polskiego społeczeństwa we wszystkich ważnych wymiarach warunków i jakości życia. Tak narodziła się Diagnoza Społeczna – nasze największe osiągnięcie badawcze. Na pierwszą edycję w 2000 r. Wiesiek załatwił pieniądze w rządzie Jerzego Buzka, obiecując dostarczyć obiektywny obraz efektów czterech reform. Obraz nie wszystkich w rządzie zachwycił. Jedna pani minister była wręcz wściekła i wyrażała sprzeciw wobec kontynuacji badania. Ale my byliśmy zdeterminowani i poszukaliśmy pieniędzy na kolejne edycje w firmach komercyjnych. Znajomości Wieska bardzo nam w tych po-

szukiwaniach pomogły. Przeszliśmy pod skrzydła organizacyjne PTS, w którym Wiesiek pełnił funkcję wiceprezesa. I tak dotarliśmy do ostatniej, ósmej edycji Diagnostyki w 2015 r. Łącznie przebadaliśmy 80 tys. Polaków, zaglądając we wszystkie zakamarki ich gospodarstw domowych i duszy.

Wiesiek był skarbnicą wiedzy i anegdot. Wyrwany w środku nocy mógł sypać danymi statystycznymi na dowolny temat. Nie trzeba było się trudzić i sięgać do raportów GUS, wystarczył telefon do niego. Miał w wielu sprawach wyrobione zdanie i potrafił rzeczowo bronić swoich racji. Był wulkanem energii. Nie mógł spokojnie przysiąc nawet na emeryturze. Ilekroć odwiedzał nas z żoną i synem, dowiadaliśmy się o jego kolejnych pomysłach na zagospodarowanie wolnego czasu.

Dlaczego tak dobrze nam się współpracowało? Odpowiedź jest prosta: był słowny, rzetelny i uczciwy, a ponadto miał poczucie humoru.

Wiesiek nauczył mnie, że statystyka ma duszę, o ile umie się ją czytać.

Janusz Czapieński

*
* *

Władysław Wiesław Łagodziński w zasadzie całe swoje życie zawodowe poświęcił statystyce. Jako młody człowiek, absolwent Uniwersytetu Warszawskiego, w 1968 r. rozpoczął pracę w GUS i pozostał z nim związany niemal do końca życia. Pełnił wiele ważnych funkcji. Był m.in. rzecznikiem prasowym GUS, dyrektorem Departamentu Informacji i radcą Prezesa GUS. Zainicjował i współtworzył blisko 250 badań ankietowych, w tym bardzo ważny cykl badań Diagnostyka Społeczna. Od 2011 r. pełnił funkcję przewodniczącego warszawskiego oddziału PTS. Jako pasjonat i propagator statystyki publicznej i edukacji statystycznej społeczeństwa zorganizował 41 seminariów popularnonaukowych, na które zapraszał znamienitych prelegentów. Mając ponad 50 lat doświadczenia w pracy w statystyce publicznej, cały czas dawał przykład swoim zapałem i aktywnością młodszemu pokoleniu. Uczestniczył w pracach komisji okręgowej Olimpiady Statystycznej. Wspierał konkursy statystyczne organizowane przez Urząd Statystyczny w Warszawie – fundował nagrody jako przewodniczący PTS i włączał się w prace komisji konkursowej. Swoją miłością do statystyki dzielił się podczas licznych konferencji i seminariów poświęconych badaniom statystycznym. Niewątpliwie można go uznać za jednego z najwybitniejszych polskich statystyków. Był członkiem honorowym PTS, wiceprezesem PTS w latach 2010–2018, kierownikiem Biura Badań i Analiz Statystycznych przy Radzie Głównej PTS, długoletnim członkiem Kolegium Redakcyjnego „Wiadomości Statystycznych”, znakomitą specjalistą i cenionym współpracownikiem, a także autorem wielu książek, artykułów, analiz i ekspertyz.

Swoją pasję зараżał innych. Był nietuzinkowy, jedyny w swoim rodzaju. W czasie wielu rozmów, jakie odbyliśmy, szczególną uwagę poświęcał Szczecinowi, do którego

miał niezwykle sentyment. Planowaliśmy seminarium na jego temat, ale niestety nie udało się nam zrealizować tego pomysłu. Był osobowością polskiej statystyki, człowiekiem z charakterem, o ugruntowanych zasadach, który potrafił słuchać argumentów. Bez reszty oddany działalności PTS i jego misji doskonalenia i popularyzacji dorobku statystyki publicznej. Będzie go nam wszystkim bardzo brakowało.

Waldemar Tarczyński

*

* *

Nasza znajomość z Wieskiem datuje się od 1983 r., kiedy oboje rozpoczęliśmy pracę w GUS. Z uwagi na niebywałą bezpośredniość Wieska zaprzyjaźniliśmy się niemalże od pierwszego kontaktu. Pracowaliśmy razem aż do 2004 r., w którym oboje odeszliśmy z GUS do Urzędu Statystycznego w Warszawie. W tym samym roku Wiesiek przeszedł – acz niechętnie – na emeryturę. Nie wyobrażał sobie, że nie będzie miał już wpływu na to, co się dzieje w statystyce publicznej, i podjął pracę w US w Warszawie, co przyjęliśmy z dużą satysfakcją.

Od pierwszego dnia bardzo angażował się w życie zawodowe urzędu. Wszędzie było go pełno, szybko nawiązał doskonały kontakt ze wszystkimi pracownikami. Jego ogromne doświadczenie, którym chętnie się dzielił, przełożyło się m.in. na doskonalenie produktów statystycznych wytwarzanych przez urząd. Owocem tego było m.in. podniesienie poziomu redakcyjnego i merytorycznego publikacji *Panorama dzielnic Warszawy*. Zamiłowanie do zgłębiania prawdziwej historii Warszawy skłoniło go do drażnienia zasobów archiwalnych US. Zaowocowało to odkryciem dawno zapomnianych dokumentów, w tym publikacji zawierających wyniki spisu ludności przeprowadzonego tuż po wyzwoleniu miasta. Chcąc, aby odnalezione dokumenty dotarły do jak największej liczby pasjonatów historii Warszawy, opracował wraz z innymi pracownikami US publikację *Warszawa Feniksem XX wieku. Statystyka Warszawy w 1945 r.*

Bardzo chętnie dzielił się swoim doświadczeniem i kompetencjami, szczególnie z adeptami statystyki, których bezinteresownie wspierał podczas ich studiów na uczelniach. Z zapałem podejmował się szkoleń ankietatorów statystycznych i przekazywał im unikalne rady socjologa dotyczące zachowań respondentów, co potem przekładało się na kompletność i lepszą jakość badań ankietowych prowadzonych przez urząd. Nie wyobrażał sobie dnia bez kontaktu z mediami. W US w Warszawie mógł nadal oddawać się swojej pasji, tym razem w wymiarze lokalnym.

Od początku swojej pracy zawodowej aktywnie uczestniczył w działalności PTS. Od 2011 r. pełnił funkcję przewodniczącego Rady Oddziału Warszawskiego PTS. To właśnie dzięki Wieskowi w tym okresie oddział warszawski przeżył swoją drugą młodość. Wiesiek, odkąd objął kierownictwo, stawiał na (zaniedbaną według niego) popularyzację wiedzy o statystyce publicznej wśród licznej rzeszy odbiorców.

Zaowocowało to m.in. zorganizowaniem cyklu kilkudziesięciu seminariów naukowych. Dzięki niemu pracownicy statystyki mieli niepowtarzalną możliwość wysłuchania wielu autorytetów z zakresu m.in. socjologii, demografii i gospodarki oraz uczestniczenia w dyskusji z nimi.

Właściwie do ostatnich dni był z nami w kontakcie. Z wielkim żalem przyjęliśmy wiadomość o jego odejściu. W urzędzie nigdy już nie będzie tak jak przez lata jego obecności. Odszedł człowiek wyjątkowy, jeśli chodzi o zapał do propagowania w społeczeństwie wiedzy o systemie statystyki publicznej.

Zofia Kozłowska
Krzysztof Kowalski

*
* *

Poznałam Wiesława Łagodzińskiego w 1991 r., kiedy rozpoczęłam pracę w ówczesnym dziale zbiorów archiwalnych i specjalnych Centralnej Biblioteki Statystycznej (dziś jest to Centralne Archiwum GUS). Wiesław zaproponował mi wówczas – w porozumieniu z moim ówczesnym kierownikiem Janem Bergerem – opracowanie akt NSZZ „Solidarność” w GUS. Chętnie tę pracę wykonałam, a Wiesław stale monitorował postęp archiwizacji i opracowywania dokumentów i przypominał mi, jaką ważną pracę wykonuję. Chętnie dzielił się ze mną także wiedzą pozaźródłową dotyczącą opracowania tych akt. Już wtedy zwróciłam uwagę na jego niezwykłą cechę – pasję, z jaką wykonywał pracę na rzecz statystyki publicznej i zachęcał do niej innych.

Wiele lat później został moim szefem – z polecenia prezesa GUS nadzorował działalność CBS, której jestem dyrektorem. Nie zapomnę nigdy jego zaangażowania w rozwój naszej placówki, zwłaszcza w jej nieustanne unowocześnianie i rozszerzanie zakresu prac informacyjnych, w tym publikacyjnych. Wiesław dosłownie żył tym, co robił. Niezapomniane były odprawy naczelników wydziałów Departamentu Informacji i dyrektorów CBS i ZWS, podczas których w wielogodzinnych dyskusjach wypracowywaliśmy pod jego kierunkiem zasady pracy, zwłaszcza dotyczące udostępniania informacji i wprowadzania w tym zakresie najwyższych standardów.

Pod koniec życia Wiesław Łagodziński przekazał swój pokaźny księgozbiór do CBS, za co jestem mu szczególnie wdzięczna.

Bożena Łazowska

*
* *

Nie pamiętam dokładnie, kiedy po raz pierwszy w swoim zawodowym życiu zetknęłam się z Wiesławem Łagodzińskim, ale ponieważ miał on bardzo wyrazistą osobowość i otaczała go nieco mityczna aura, to zanim poznałam go osobiście, znałam go już ze słyszenia. Wiedziałam, że to ktoś, kto ma duży wpływ na funkcjonowanie GUS

i urzędów statystycznych. Zawsze bardzo pewny siebie i swobodny, barwnie i z pasją opowiadał o statystyce. W roli rzecznika GUS był w razie potrzeby bardzo skuteczną zaporą przed natarczywymi mediami, ale przede wszystkim umiał zainteresować dziennikarzy tematami, jakimi żyła statystyka. Był zawsze widoczny i skupiał uwagę wszędzie tam, gdzie się pojawił. Jego tubalny głos było słychać z daleka.

Zadzwoił kiedyś, wiele lat temu, do urzędu. Ówczesny dyrektor był nieobecny, więc mnie przyszło podjąć rozmowę z Wiesławem. Pamiętam, że byłam zaskoczona gwałtownym tonem jego słów, ale ta gwałtowność stopniowo malała i zmieniała się w zaciekawienie – bo ku jego zaskoczeniu starałam się przebić przez potok wypowiedzianych szybko poleceń dla dyrektora i dopytać o kontekst. Był to swego rodzaju bohaterski wyczyn, ponieważ niewielu śmiało „przerywać Łagodzińskiemu”. Wiem z późniejszych naszych kontaktów zawodowych, że właśnie wtedy mnie zapamiętał. A potem w ciągu wielu lat spotykaliśmy się przy różnych okazjach, m.in. na konferencjach naukowych czy w ramach prac PTS lub też po prostu na korytarzach GUS.

Halina Woźniak

*

* *

GUS, IV piętro, korytarz C, 2000 r. (początek mojej pracy w urzędzie). Odwiedziła mnie siostra. Rozmawiamy na korytarzu. Przechodzi pan Wiesław Łagodziński.

WŁ: Kto to jest, Alusiu?

Ja: To moja młodsza siostra Ania, panie dyrektorze.

WŁ: A taka niepodobna?

Ja: Ja jestem podobna do mamusi, a ona do tatusia.

WŁ: Hmm, to ona miała więcej szczęścia.

Moja mama uwielbia tę anegdotkę.

Alicja Koszela

*

* *

Trzymam w ręku medal wydany w 2008 r. z okazji 90-lecia GUS. Od tego czasu minęło już kilkanaście lat, ale we mnie wciąż żyje wspomnienie pamiętnej uroczystości. Byłem wtedy dyrektorem administracyjnym w GUS, więc znalazłem się w grupie osób przygotowujących obchody. Najpierw była msza św. w kościele akademickim św. Anny, a następnie uroczystość na Zamku Królewskim w Warszawie i sesja naukowa w Bibliotece Narodowej.

Dużo się tego dnia działo, ale dla mnie najważniejsza była organizacja uroczystości na Zamku, gdzie w sali balowej odbyła się część oficjalna, z zaproszonymi gośćmi, przemowami itd., a prowadził ją nasz Władysław. Tylko on mógł temu podołać, było to bowiem ogromnie trudne zadanie. Przygotowano scenariusz, ale nie przeprowa-

dzono próby generalnej. Mało tego, sytuacja ciągle się zmieniała, trzeba było improwizować, bo do końca nie było wiadomo, kto z zaproszonych gości przyjdzie, a kto wystosuje list. Władysław osobiście witał gości, zapoznawał ich z planem spotkania i zapraszał do zabrania głosu. Miałem wrażenie, jakby robił to zawodowy konferansjer. Wiedziałem, że wypadnie dobrze, że się nie pogubi, a jeżeli nawet, to skwituje to jakimś dowcipem – w swoim stylu, błyskotliwie i lekko.

Po tej uroczystości prowadził jeszcze sesję w Bibliotece Narodowej. Przez cały czas był w swoim żywiole.

Maciej Adamowicz

*

* *

Wiesław żył statystyką i dla statystyki, ale nasyconej myśleniem socjologicznym, które emanowało z niego – w moim odczuciu – w każdej z ról, w jakich występował; uosabiał socjologizację statystyki. Miałem, z mniejszymi lub większymi przerwami, możliwość obserwowania tej wyróżniającej go aktywności w GUS od czasów prezesa Wincentego Kawalca. Był w swoim żywiole, gdy w grę wchodziło profesjonalne mówienie o liczbach i o ludziach za pomocą liczb. Mówił o nich jak Ciszewski o piłce – z pasją i koncentracją na tej drugiej, niewidzialnej dla większości stronie. Nasycił swoje narracje emocjonalnie i czasami bardziej niż na komunikowaniu, nawet w mediach, zależało mu na przekonaniu audytorium o wadze przekazywanej informacji.

Z prawdziwą satysfakcją obserwowałem pozytywne efekty tego typu przekazu, gdy w połowie lat 90. Wiesław jako ekspert w programie Banku Światowego prowadził na Łotwie i Litwie szkolenia lokalnych ankietatorów. Program był poświęcony budowie nowoczesnego systemu statystyki gospodarstw domowych. Dzięki swojej bezpośredniości i sprawności językowej (na początku on jeden miał odwagę mówić publicznie po rosyjsku) Wiesław szybciej i lepiej od innych szkoleniowców, wspieranych przez tłumaczy, stwarzał atmosferę swojskości i zaufania, tak cenną w tego typu misjach.

Chociaż żył statystyką w sposób widoczny dla wszystkich, to miałem okazję się przekonać, że jego najprawdziwsze oddanie statystyce wyrażało się w działaniach zupełnie nierzucających się w oczy, a czasami wręcz czysto wewnętrznych. Jak np. w 2009 r., gdy dzięki jego staraniom mogłem uczestniczyć w 59. Światowym Kongresie Statystyki (World Statistics Congress) w Durbanie (RPA). „Zrobię ci to! – zapewniał mnie Wiesław jako sekretarz PTS. – Musimy być w programie WSC”. Czerpał prawdziwą satysfakcję z możliwości zrobienia czegoś dla innych. Na wspieranie przez niego obecności i widoczności statystyki polskiej za granicą mogliśmy liczyć także przy innych okazjach, np. w odniesieniu do „Statistics in Transition new se-

ries”. Zrobił wiele dla życia naukowego środowiska statystycznego w Polsce – na pewno nie byłoby ono takie samo bez jego nieocenionego udziału i wkładu.

Włodzimierz Okrasa

*
* *

Po raz pierwszy spotkałem Wiesława w 2000 r. na wspólnym posiedzeniu Rady Programowej i Kolegium Redakcyjnego „Wiadomości Statystycznych”. Oczywiście znałem go wcześniej – bo znali go właściwie wszyscy – jako charyzmatycznego rzecznika prasowego GUS. Donośny, radiowy tembr głosu i specyficzny sposób bycia wyróżniały go nie tylko w mediach, lecz także na naszych resortowych naradach. Odniosłem wrażenie, że już podczas pierwszego spotkania udało się nam nawiązać nić sympatii. Pamiętam doskonale, że zaprosił mnie do redakcji „WS”, abym zobaczył, jak wygląda w praktyce ocena i wybór artykułów do publikacji.

Od 2005 r. spotykaliśmy się już systematycznie na comiesięcznych posiedzeniach Kolegium Redakcyjnego, z których zapamiętałem go jako bardzo aktywnego uczestnika. Odkąd zostałem redaktorem naczelnym, wielokrotnie prowadziliśmy dyskusje na temat kierunku rozwoju „WS”. Często mieliśmy odmienne zdania. Wiesław podkreślał, że nasz miesięcznik nie powinien przekształcać się w tak szybkim tempie w czasopismo typowo akademickie.

Nasza znajomość nie ograniczała się do relacji czysto zawodowych. Braliśmy wspólnie udział w spotkaniach integracyjnych ze statystykami z Ukrainy i ze Słowacji, zwiedzaliśmy piękne słowackie zamki, wędrowaliśmy po Bieszczadach. Wiesiek był szczególnie dumny ze zdobycia Połoniny Caryńskiej podczas swojego ostatniego wyjazdu w góry. Był człowiekiem pełnym zapału i energii. Potrafił zaskakiwać. Rozmowy z nim nigdy nie były nudne, a jego opinie, czasem celowo prowokacyjne, pozwalały spojrzeć na sprawę z innej perspektywy.

Marek Cierpiął-Wolan

Ważniejsze publikacje

Artykuły w czasopismach

Problemy analizy przestrzennego rozmieszczenia instytucji i urzędów kulturalnych. (1974). *Statystyk Terenowy*, (3), 18–19.

Uczestnictwo w kulturze [1972 r.] (wstępne wyniki i badania). (1974). *Wiadomości Statystyczne*, (1), 29–32.

Rozwój statystycznych badań kultury w GUS w latach 1946–1974. (1975). *Wiadomości Statystyczne*, (1), 38–39.

Upowszechnianie kultury wśród dzieci i młodzieży. (1979). *Wiadomości Statystyczne*, (5), 43–47.

Problemy metodologii badań ruchu turystycznego. (1980). *Wiadomości Statystyczne*, (11), 13–15 i (12), 22–24.

- Problemy badań uczestnictwa w kulturze. (1987). *Wiadomości Statystyczne*, (8), 8–11.
- Modułowe badania ankietowe w Zintegrowanym Systemie Badań Gospodarstw Domowych. (1991). *Wiadomości Statystyczne*, (5), 13–16.
- Statystyka społeczna w magazynie „Dzień dobry” – informacja, analiza, ocena, popularyzacja. (1992). *Wiadomości Statystyczne*, (4), 41–42.
- Podstawowe zasady statystyki oficjalnej oraz ich prawne i etyczne aspekty w krajach okresu przejściowego. (1993). *Wiadomości Statystyczne*, (12), 36–37.
- Pomoc społeczna dla rodzin z dziećmi do 13 lat. (1993). *Wiadomości Statystyczne*, (6), 15–16.
- Zintegrowany System Badań Warunków Życia Ludności w latach 1993–1995. (1993). *Wiadomości Statystyczne*, (6), 8–12.
- The challenges of the Polish statistics – public lectures of the problems and conditions of the transformation of the Polish statistics. (1994). *Statistics in Transition*, 1(3), 395–399.
- Aktywna polityka informacyjna. (1995). *Wiadomości Statystyczne*, (12), 1–6.
- Dostęp do informacji a istota tajemnicy statystycznej. (1995). *Wiadomości Statystyczne*, (11), 20–24.
- Doświadczenia Głównego Urzędu Statystycznego w obsłudze klientów. (1995). *Gospodarka Narodowa*, (11/12), 61–63.
- Prawo do informacji statystycznej. (1995). *Rzeczpospolita* (wyd. zasadnicze), (152), 20–22.
- Przyczynek do przestrzennego zróżnicowania samobójstw. (1995). *Wiadomości Statystyczne*, (9), 45–47.
- Dzieci w obszarach pogranicza. (1999). *Wiadomości Statystyczne*, (3), 29–32.
- Niektóre problemy badań sytuacji gospodarstw domowych na pograniczu Polski. (1999). *Wiadomości Statystyczne*, (11), 40–43.
- [Współautor: J. Czapiński]. Podstawowe kwestie metodologii badań sytuacji dzieci w obszarach przygranicznych. (1999). *Wiadomości Statystyczne*, (9), 52–55.
- The relationship between official statistics and the media in Poland: the role of information service and the confrontation and cooperation zones. (2009). *Statistics in Transition new series*, 10(1), 155–162.
- Diagnoza Społeczna 2000–2013. (2014). *VIP Magazyn* (wydanie specjalne), 25–26.
- Finanse Polskiego Towarzystwa Statystycznego. Biuro Badań i Analiz Statystycznych przy Radzie Głównej Polskiego Towarzystwa Statystycznego. (2014). *VIP Magazyn* (wydanie specjalne), 23–24.
- [Współautor: B. Łazowska]. 60 lat „Wiadomości Statystycznych”. (2017). *Wiadomości Statystyczne*, (4), 5–15.

Publikacje w wydawnictwach zbiorowych

- Klasyfikacyjne i definicyjne problemy stosowania mierników w statystycznych badaniach działalności kulturalnej. (1976). W: A. Machnowski (oprac.), *Zagadnienia metodologiczne statystyki społeczno-demograficznej: konferencja naukowa, Warszawa 25 IV 1975 r.* (s. 88–103). Warszawa: Główny Urząd Statystyczny.
- Uczestnictwo w kulturze uczniów i studentów. (1978). W: W. Kondrat (oprac. merytoryczne), *Rola młodzieży w życiu społeczno-gospodarczym kraju: konferencja naukowa, Gdańsk 10–11 V 1976 r.* (s. 261–273). Warszawa: Główny Urząd Statystyczny.
- Turystyka i wypoczynek. (1986). W: W. Sadowski (red.), *Warunki życia ludności w latach 1981–1985* (s. 116–117). Warszawa: Główny Urząd Statystyczny.
- Ważniejsze dane o upowszechnianiu kultury. (1986). W: W. Sadowski (red.), *Warunki życia ludności w latach 1981–1985* (s. 106–111). Warszawa: Główny Urząd Statystyczny.
- Badanie dostępu i uczestnictwa w kulturze w Zintegrowanym Systemie Badań Gospodarstw Domowych. (1987). W: S. Jońca, S. Mierzejewski (oprac. merytoryczne i red.), *Problemy integracji statystycznych badań gospodarstw domowych* (s. 168–167). Warszawa: Główny Urząd Statystyczny.

- Warunki bytu a potrzeby młodzieży. (1991). W: W. Łagodziński (oprac. merytoryczne), *Jakość życia i warunki bytu* (s. 89–107). Warszawa: Główny Urząd Statystyczny, Polskie Towarzystwo Statystyczne.
- Biuro Badań i Analiz Statystycznych przy RG PTS. (1992). W: Polskie Towarzystwo Statystyczne, *Polskie Towarzystwo Statystyczne 1912–1992* (s. 123–139). Warszawa.
- Dekalog zasad statystyk oficjalnych jako przesłanki do kształtowania się nowego ładu informatycznego. (1995). W: S. Jońca (oprac. red.), *Obchody jubileuszowe 200-lecia statystyki polskiej i 75-lecia Głównego Urzędu Statystycznego: sesja naukowo-historyczna, Zamek Królewski w Warszawie, 12 lipca 1993 r.* (s. 75–78). Warszawa: Główny Urząd Statystyczny, Polskie Towarzystwo Statystyczne.
- Ewolucja badań warunków życia ludności i jej odniesienia do transformacji społeczno-ekonomicznej w Polsce. (1995). W: S. Jońca (oprac. red.), *Obchody jubileuszowe 200-lecia statystyki polskiej i 75-lecia Głównego Urzędu Statystycznego: sesja naukowo-historyczna, Zamek Królewski w Warszawie, 12 lipca 1993 r.* (s. 124–131). Warszawa: Główny Urząd Statystyczny, Polskie Towarzystwo Statystyczne.
- BBiAS – Biuro Badań i Analiz Statystycznych przy Radzie Głównej PTS. (2012). W: K. Kruska (red.), *Polskie Towarzystwo Statystyczne 1912–2012* (s. 157–168). Warszawa: Polskie Towarzystwo Statystyczne.

Wydawnictwa zwarte

- [Oprac. pod kier. L. Adamczuka; autorzy koncepcji badania: W. Łagodziński, L. Adamczuk, M. Strąk]. (1981). *Uczestnictwo ludności w kulturze 1979. Raport z badania statystyczno-socjologicznego*. Warszawa: Główny Urząd Statystyczny.
- [Oprac. W. Łagodziński]. (1987). *Uczestnictwo w kulturze w 1985 r.* Warszawa: Główny Urząd Statystyczny.
- [Oprac. W. Łagodziński]. (1987). *Uczestnictwo w turystyce ludności Polski w 1986 r.* Warszawa: Główny Urząd Statystyczny.
- Warunki życia i potrzeby młodzieży*. (1989). Warszawa: Główny Urząd Statystyczny.
- [Oprac. merytoryczne W. Łagodziński]. (1991). *Jakość życia i warunki bytu*. Warszawa: Główny Urząd Statystyczny, Polskie Towarzystwo Statystyczne.
- Uczestnictwo w kulturze: podstawowe wyniki badań reprezentacyjnych z lat: 1972, 1979, 1985, 1988 i 1990*. (1992). Warszawa: Główny Urząd Statystyczny.
- [Oprac. J. Berger, H. Ducal, W. Orlik pod kier. W. Łagodzińskiego]. (2002). *213 lat spisów ludności w Polsce 1789–2002*. Warszawa: Główny Urząd Statystyczny.
- [Red. zespół pod kier. W. Łagodzińskiego]. (2003). *85 lat Głównego Urzędu Statystycznego (1918–2003): przeszłość dla przyszłości*. Warszawa: Główny Urząd Statystyczny.
- Szanse i zagrożenia uczestnictwa w kulturze w latach 1990–2003: w świetle wyników badań Głównego Urzędu Statystycznego*. (2004). Warszawa: Narodowe Centrum Kultury.
- [Oprac. merytoryczne i red. W. W. Łagodziński]. (2008). *Jubileusz 90-lecia Głównego Urzędu Statystycznego 1918–2008: uroczyste obchody pod honorowym patronatem Prezydenta Rzeczypospolitej Polskiej Lecha Kaczyńskiego na Zamku Królewskim w Warszawie, 11 lipca 2008 r. Seminarium naukowo-historyczne „Kluczowe problemy polskiej statystyki” w Bibliotece Narodowej w Warszawie, 11 lipca 2008 r.* Warszawa: Główny Urząd Statystyczny.
- [Oprac. J. Karczmarski; red. W. W. Łagodziński]. (2008). *Tajemnica statystyczna w badaniach statystyki publicznej: istota problemu*. Warszawa: Główny Urząd Statystyczny.
- [Pod kier. nauk. J. Czapieńskiego; oprac. raportu W. W. Łagodziński i in.]. (2010). *Badanie warunków i jakości życia oraz zachowań ekonomicznych w gospodarstwach domowych województwa podkarpackiego*. Warszawa, Rzeszów: Urząd Statystyczny w Rzeszowie, Polskie Towarzystwo Statystyczne.
- [Zespół aut.: W. W. Łagodziński i in.]. (2012). *Warszawa Feniksem XX wieku. Statystyka Warszawy w 1945 r.* [CD-ROM]. Warszawa: Urząd Statystyczny w Warszawie, Polskie Towarzystwo Statystyczne.

- [Oprac. A. Brzezińska-Skarżyńska, K. Madoń-Mitzner; aut. tekstów: B. Czerwińska-Jędrusiak, W. W. Łagodziński, J. Mencwel, J. Pałka, T. Stempowski]. (2016). *Na nowo: warszawiacy 1945–1955* (tłum. B. Herchenreder). Warszawa: Dom Spotkań z Historią.
- [Zespół red.: przewodn. i red. gł. F. Kubiczek, członkowie: W. Adamczewski, H. Dmochowska, C. Kukło, C. Leszczyńska, W. W. Łagodziński, B. Łazowska, J. Łukasiewicz, G. Szydłowska; współpraca merytoryczna P. Ciecieląg]. (2017). *Historia Polski w liczbach: t. 4. Statystyka Polski, dawniej i dzisiaj*. Warszawa: Główny Urząd Statystyczny.
- [Zespół red.: przewodn. i red. gł. F. Kubiczek, członkowie: W. Adamczewski, H. Dmochowska, C. Kukło, C. Leszczyńska, W. W. Łagodziński, B. Łazowska, J. Łukasiewicz, G. Szydłowska]. (2018). *Historia Polski w liczbach: t. 5. Polska 1918–2018*. Warszawa: Główny Urząd Statystyczny.

Remigiusz Tworzydło

WYDAWNICTWA GUS. LUTY 2021

PUBLICATIONS OF STATISTICS POLAND. FEBRUARY 2021

W ofercie wydawniczej Głównego Urzędu Statystycznego z ubiegłego miesiąca warto zwrócić uwagę na następujące publikacje:

Among Statistics Poland's last month's publications, we would like to recommend:



Tytuł: *Pokolenie gniazdowników w Polsce*

Title: *Generation of young adults living with their parents in Poland*

Język: polski (przedmowa, spis treści, wstęp oraz synteza dodatkowo w języku angielskim)

Language: Polish (preface, contents, introduction and executive summary additionally in English)

Dodatkowe informacje: opracowanie dostępne w wersji elektronicznej

Additional information: publication available in the electronic version

Praca eksperymetalna nakreślająca całościowy portret pokolenia gniazdowników w Polsce. Opiera się na danych z rejestrów administracyjnych i na wynikach wcześniejszych prac badaw-

czych dotyczących procesu łączenia mieszkańców kraju w rodziny oraz aktywności ekonomicznej ludności. Na ich podstawie określono zróżnicowanie przestrzenne poziomu gniazdownictwa i dokonano szczegółowej charakterystyki młodych Polaków mieszkających z rodzicami (m.in. według wieku, płci, aktywności ekonomicznej i edukacyjnej oraz średnich miesięcznych dochodów w 2018 r.). Ustalono również zależność poziomu gniazdownictwa od modelu rodziny i aktywności zawodowej rodziców. Wnioski płynące z opracowania mogą być pomocne w formułowaniu rekomendacji m.in. dla polityki rodzinnej, mieszkaniowej, edukacyjnej czy rynku pracy w zakresie tworzenia warunków sprzyjających uzyskiwaniu samodzielności ekonomicznej i mieszkaniowej przez młode pokolenie.



Tytuł: Jakość życia osób starszych w Polsce

Title: Quality of life of elderly people in Poland

Język: polski (przedmowa, spis treści oraz synteza dodatkowo w języku angielskim)

Language: Polish (preface, contents and executive summary additionally in English)

Dodatkowe informacje: opracowanie dostępne w wersji elektronicznej

Additional information: publication available in the electronic version

Opracowanie analityczne wpisujące się w serię publikacji na temat różnych aspektów jakości życia osób starszych w Polsce. Omówiono w nim zagadnienia dotyczące sytuacji materialnej, aktywności w czasie wolnym, korzystania z internetu,

relacji społecznych, wyznawanych wartości i życia religijnego. Osobny rozdział poświęcono subiektywnemu dobrobytowi seniorów, z uwzględnieniem oceny stanu emocjonalnego i ogólnego zadowolenia z życia. Opracowanie zawiera także syntetyczną analizę porównawczą jakości życia osób starszych w Polsce i Unii Europejskiej. Przedstawione dane odnoszą się do okresu sprzed pandemii Covid-19 i mogą stanowić punkt wyjścia do dalszych badań i analiz obrazujących wpływ pandemii zarówno na obiektywne warunki życia, jak i na subiektywny dobrobyt osób starszych.

W lutym br. ukazały się ponadto:

- *Badanie Aktywności Ekonomicznej Ludności (folder dla rodzin biorących udział w badaniu aktywności ekonomicznej ludności) – III kwartał 2020 r.;*
- „Biuletyn statystyczny” nr 1/2020;
- *Ceny robót budowlano-montażowych i obiektów budowlanych (grudzień 2020 r.);*
- *Dochody i warunki życia ludności Polski – raport z badania EU-SILC 2019;*
- *Działalność innowacyjna przedsiębiorstw w latach 2017–2019;*
- *Działalność przedsiębiorstw niefinansowych w 2019 r.;*
- *Koniunktura w przetwórstwie przemysłowym, budownictwie, handlu i usługach 2000–2021 (luty 2021);*
- *Produkcja ważniejszych wyrobów przemysłowych w styczniu 2021 r.;*
- *Przedsiębiorstwa niefinansowe powstałe w latach 2015–2019;*
- *Sytuacja społeczno-gospodarcza kraju w styczniu 2021 r.;*
- „Wiadomości Statystyczne. The Polish Statistician” nr 2/2021;
- *Zeszyt metodologiczny. Badanie kondycji gospodarstw domowych (postaw konsumenckich);*

- *Zeszyt metodologiczny. Zapotrzebowanie rynku pracy na pracowników według zawodów;*
- *Zdrowie i ochrona zdrowia w 2019 r.*

Wersje elektroniczne wszystkich publikacji GUS są dostępne na stronie stat.gov.pl/publikacje/publikacje-a-z.

Electronic versions of all the publications by Statistics Poland are available at stat.gov.pl/en/publications.

Justyna Gustyn

Główny Urząd Statystyczny, Departament Opracowań Statystycznych
Statistics Poland, Statistical Products Department

DLA AUTORÓW FOR THE AUTHORS

(for the English translation of the information given below, please visit ws.stat.gov.pl/ForAuthors)

W „Wiadomościach Statystycznych. The Polish Statistician” („WS”) zamieszczane są artykuły o charakterze naukowym poświęcone teorii i praktyce statystycznej, które prezentują wyniki oryginalnych badań teoretycznych lub analitycznych wykorzystujących metody statystyki matematycznej, opisowej bądź ekonometrii. Ukazują się również artykuły przeglądowe, recenzje publikacji naukowych oraz inne opracowania informacyjne. W czasopiśmie publikowane są prace w języku polskim i angielskim.

Od 2007 r. „WS” znajdują się na liście polskich punktowanych czasopism naukowych MEiN. Zgodnie z komunikatem Ministra Edukacji i Nauki z dnia 9 lutego 2021 r. w sprawie wykazu czasopism naukowych i recenzowanych materiałów z konferencji międzynarodowych „WS” otrzymały 40 punktów.

„Wiadomości Statystyczne. The Polish Statistician” są udostępniane w następujących bazach indeksacyjnych i repozytoriach: Agro, BazEkon, Central and Eastern European Academic Source (CEEAS), Central and Eastern European Online Library (CEEOL), Central European Journal of Social Sciences and Humanities (CEJSH), EBSCO Discovery Service, European Reference Index for the Humanities and Social Sciences (ERIH Plus), ICI Journals Master List, ICI World of Journals, Norwegian Register for Scientific Journals and Publishers (The Nordic List) oraz POL-index.

Za publikację artykułów na łamach „WS” autorzy nie otrzymują honorariów ani nie wnoszą opłat.

1. Zgłaszanie artykułów

Prace należy przysyłać na adres: redakcja.ws@stat.gov.pl.

Artykuł powinien być utrzymany w formie bezosobowej i zawierać streszczenie, słowa kluczowe, kod/kody JEL oraz ORCID, adres e-mail i afiliację autora. Tytuł, streszczenie i słowa kluczowe powinny być podane w językach polskim i angielskim.

Jeżeli w pracy występują tablice, wykresy lub mapy, powinny być umieszczone w treści artykułu. W osobnym pliku (najlepiej w formacie Excel) należy podać dane do wykresów.

Autor jest zobowiązany do podania w artykule wszelkich źródeł finansowania badań będących podstawą pracy. Jeżeli doszło do zaprezentowania podobnych materiałów podczas konferencji lub sympozjum naukowego, to podczas składania tekstu do publikacji w „WS” należy poinformować o tym redakcję.

Prosimy o niestosowanie stylów i ograniczenie formatowania do wymogów redakcyjnych. Więcej informacji w rozdziale *Wymogi redakcyjne*.

Razem z artykułem należy przesłać skan oświadczenia (do pobrania ze strony internetowej czasopisma) o oryginalności pracy i niezłożeniu jej w innym wydawnictwie, zawierającego zgodę na przeniesienie autorskich praw majątkowych, numer ORCID, afiliację lub afiliacje oraz dane kontaktowe autora, wraz ze wskazaniem proponowanego działu czasopisma. Oryginał oświadczenia należy wysłać na adres: Redakcja „Wiadomości Statystycznych. The Polish Statistician”, Główny Urząd Statystyczny, al. Niepodległości 208, 00-925 Warszawa.

Załączenie skanu oświadczenia jest warunkiem poddania pracy ocenie wstępnej i skierowania do recenzji.

2. Przebieg prac redakcyjnych

Zgłoszony artykuł jest oceniany i opracowywany w czteroetapowym procesie:

1. **Ocena wstępna**, dokonywana przez redakcję. Polega na weryfikacji naukowego charakteru artykułu oraz jego struktury i zawartości pod kątem wymogów redakcyjnych, a także zgodności tematyki z profilem czasopisma. Autor uzupełnia i poprawia artykuł stosownie do uwag redakcji, a w przypadku nieuwzględnienia danej uwagi uzasadnia swoje stanowisko. **Razem z poprawionym artykułem autor przesyła w osobnym pliku zanonimizowaną wersję pracy, przeznaczoną do recenzji.** Anonimizacja polega na utajnieniu nazwiska autora (także we właściwościach pliku), usunięciu podziękowań i informacji o źródłach finansowania, a także innych informacji wskazujących na afiliację lub umożliwiających zidentyfikowanie autora. Warunkiem skierowania pracy do recenzji jest potwierdzenie oryginalności tekstu uzyskane za pomocą systemu antyplagiatowego Similarity Check. W przypadku wykrycia znacznego podobieństwa do innych prac artykuł zostanie odrzucony.
2. **Ocena recenzentów**, dokonywana przez specjalistów w danej dziedzinie. Artykuł oceniają dwaj recenzenci spoza jednostki naukowej, przy której afiliowany jest autor; w przypadku pracy w języku angielskim co najmniej jeden recenzent jest afiliowany przy jednostce zagranicznej. W razie sprzecznych opinii dwóch recenzentów powoływany jest trzeci recenzent. Recenzenci kierują się kryteriami oryginalności i jakości opracowania zarówno w odniesieniu do treści, jak i formy.

Autorzy artykułów, które otrzymały pozytywne oceny, wprowadzają poprawki zalecane przez recenzentów i przesyłają do redakcji zmodyfikowaną wersję pracy. Jeśli pojawi się różnica zdań dotycząca zasadności proponowanych zmian, autorzy są zobligowani do uzasadnienia swojego stanowiska.

3. **Ocena Kolegium Redakcyjnego (KR)**, decydująca o przyjęciu pracy do publikacji. Jest dokonywana na podstawie recenzji, z uwzględnieniem opinii redaktorów tematycznego i merytorycznego. Polega m.in. na weryfikacji dokonania przez autora zmian w artykule stosownie do uwag recenzentów. KR ocenia artykuł pod względem poprawności i spójności merytorycznej oraz zaleca autorowi wprowadzenie poprawek, jeśli są one konieczne, aby praca spełniała wymogi czasopisma. Autorowi przysługuje prawo do odwołania od decyzji o niepublikowaniu artykułu. W takim przypadku powinien on skontaktować się z redakcją „WS” i przedstawić uzasadnienie. Ostateczna decyzja w tej sprawie należy do redaktora naczelnego.

W „WS” publikowane są wyłącznie te artykuły, które otrzymają pozytywną ocenę na każdym z wymienionych etapów i zostaną poprawione przez autora zgodnie z otrzymanymi uwagami (chyba że autor przedstawi argumenty uzasadniające nieuwzględnienie danej uwagi).

Artykuły przyjęte przez KR do publikacji są zamieszczane na stronie internetowej czasopisma w zakładce Early View. Znajdują się tam do czasu opublikowania w konkretnym wydaniu „WS”.

4. **Opracowanie redakcyjne, autoryzacja i korekta.** Artykuł zakwalifikowany do druku jest poddawany opracowaniu merytorycznemu i językowemu. Redakcja zastrzega sobie prawo

do zmiany tytułu i śródtytułów, modyfikowania tablic, wykresów i innych elementów graficznych oraz przeredagowania treści bez naruszenia zasadniczej myśli autora.

Po opracowaniu redakcyjnym artykuł jest przesyłany do autoryzacji. Tekst zatwierdzony przez autora, po składzie i łamaniu, jest poddawany korekcie i rewizji (II korekcie). Autor dokonuje korekty autorskiej tekstu na etapie rewizji. Wykresy i inne materiały graficzne są opracowywane na podstawie danych przekazanych przez autora i poddawane korekcie i rewizji. Autor dokonuje ich akceptacji na etapie rewizji.

W przypadku odkrycia błędów w opublikowanym artykule zamieszcza się na łamach „WS” sprostowanie, a artykuł w wersji elektronicznej jest poprawiany i umieszczany na stronie internetowej „WS” ze stosownym wyjaśnieniem.

3. Zasady etyki publikacyjnej COPE

Redakcja „WS” podejmuje wszelkie starania w celu utrzymania najwyższych standardów etycznych zgodnie z wytycznymi Komitetu ds. Etyki Publikacyjnej (COPE), dostępnymi na stronie internetowej www.publicationethics.org, oraz wykorzystuje wszystkie możliwe środki mające na celu zapobieżenie nadużyciom i nierzetelności autorskiej. Przyjęte zasady postępowania obowiązują autorów, Radę Naukową, Kolegium Redakcyjne, redakcję, pracowników Wydziału Czasopism Naukowych, recenzentów i wydawcę.

3.1. Odpowiedzialność autorów

1. Artykuły naukowe kierowane do opublikowania w „WS” powinny zawierać precyzyjny opis badanych zjawisk i stosowanych metod oraz autorskie wnioski i sugestie dotyczące rozwoju badań i analiz statystycznych. Autorzy powinni wyraźnie określić cel artykułu oraz jasno przedstawić wyniki przeprowadzonej analizy. Prezentacja efektów badań statystycznych zaprojektowanych i przeprowadzonych przez autorów wymaga opisanego zastosowania w nich metodologii. W przypadku nowatorskich metod analizy pożądanym jest podanie przykładu ilustrującego ich zastosowanie w praktyce statystycznej. Autorzy ponoszą odpowiedzialność za treści prezentowane w artykułach. W razie zgłaszania przez czytelników zastrzeżeń odnoszących się do tych treści autorzy są zobligowani do udzielenia odpowiedzi za pośrednictwem redakcji.
2. Na autorach spoczywa obowiązek zapewnienia pełnej oryginalności przedłożonych prac. Redakcja nie toleruje przejawów nierzetelności naukowej autorów, takich jak:
 - duplikowanie publikacji – ponowne publikowanie własnego utworu lub jego części;
 - plagiat – przywłaszczenie cudzego utworu lub jego fragmentu bez podania informacji o źródle;
 - fabrykowanie danych – oparcie pracy naukowej na nieprawdziwych wynikach badań;
 - autorstwo widmo (*ghost authorship*) – nieujawnianie współautorów, mimo że wniesli oni istotny wkład w powstanie artykułu;
 - autorstwo gościnne (*guest authorship*) – podawanie jako współautorów osób o znikomym udziale lub niebiorących udziału w opracowywaniu artykułu;
 - autorstwo grzecznościowe (*gift authorship*) – podawanie jako współautorów osób, których wkład jest oparty jedynie na słabym powiązaniu z badaniem.

Autorzy deklarują w stosownym oświadczeniu, że zgłaszany artykuł nie narusza praw autorskich osób trzecich, nie był dotychczas publikowany i nie został złożony w innym wydawnictwie oraz że jest ich oryginalnym dziełem, i określają swój wkład w opracowanie artykułu. Jeżeli doszło do zaprezentowania podobnych materiałów podczas konferencji lub sympozjum naukowego, to podczas składania tekstu do publikacji w „WS” autorzy są zobowiązani poinformować o tym redakcję.

3. Autorzy są zobowiązani do podania w treści artykułu wszelkich źródeł finansowania badań będących podstawą pracy.
4. Główną odpowiedzialność za rzetelność przekazanych informacji, łącznie z informacją na temat wkładu poszczególnych współautorów w powstanie artykułu, ponosi zgłaszający artykuł.
5. Autorzy zgłaszający artykuły do publikacji w „WS” biorą udział w procesie recenzji double-blind peer review, dokonywanej przez co najmniej dwóch niezależnych ekspertów z danej dziedziny. Po otrzymaniu pozytywnych recenzji autorzy wprowadzają zalecane przez recenzentów poprawki i dostarczają redakcji zaktualizowaną wersję opracowania wraz z pisemnym poświadczeniem uwzględnienia poprawek. Jeśli pojawi się różnica zdań co do zasadności proponowanych zmian, należy wyjaśnić, które poprawki zostały uwzględnione, a w przypadku ich nieuwzględnienia – uzasadnić swoje stanowisko.
6. Jeżeli autorzy odkryją w swoim maszynopisie lub tekście już opublikowanym błędy, nieścisłości bądź niewłaściwe dane, powinni niezwłocznie poinformować o tym redakcję w celu dokonania korekty, wycofania tekstu lub zamieszczenia sprostowania. W przypadku korekty artykułu już opublikowanego jego nowa wersja jest zamieszczana na stronie internetowej „WS” wraz ze stosownym wyjaśnieniem.

3.2. Odpowiedzialność Rady Naukowej, Kolegium Redakcyjnego i Wydziału Czasopism Naukowych GUS

1. Rada Naukowa (RN) kształtuje profil programowy czasopisma, określa kierunki jego rozwoju i konsultuje jego zakres merytoryczny.
2. Kolegium Redakcyjne (KR) podejmuje decyzję o publikacji danego artykułu z uwzględnieniem ocen recenzentów oraz opinii zespołu redakcyjnego. W swoich rozstrzygnięciach członkowie KR kierują się kryteriami merytorycznej oceny wartości artykułu, jego oryginalności i jasności przekazu, a także ścisłego związku z celem i zakresem tematycznym „WS”. Oceniają artykuły niezależnie od płci, rasy, pochodzenia etnicznego, narodowości, religii, wyznania, światopoglądu, niepełnosprawności, wieku lub orientacji seksualnej ich autorów.
3. Zespół redakcyjny, wyodrębniony z KR, tworzą redaktor naczelny i jego zastępca, redaktorzy tematyczni i redaktor merytoryczny. Członkowie zespołu redakcyjnego weryfikują nadane artykuły pod względem merytorycznym, oceniają ich zgodność z celem i zakresem tematycznym „WS” oraz sprawdzają spełnienie wymogów redakcyjnych i przestrzeganie zasad rzetelności naukowej. Ponadto wybierają recenzentów w taki sposób, aby nie wystąpił konflikt interesów, i dbają o zapewnienie uczciwego, bezstronnego i terminowego procesu recenzowania.
4. Za sprawny przebieg procesu wydawniczego, informowanie wszystkich jego uczestników o konieczności przestrzegania obowiązujących zasad i przygotowanie artykułów do

publikacji odpowiadają pracownicy Wydziału Czasopism Naukowych (WCN) GUS. W celu uzyskania obiektywnej oceny oryginalności nadsyłanych artykułów przed skierowaniem ich do recenzji WCN wykorzystuje system antyplagiatowy. Informacje dotyczące artykułu mogą być przekazywane przez WCN wyłącznie autorom, recenzentom, członkom RN i KR oraz wydawcy.

5. Zmiany dokonane w tekście na etapie przygotowania artykułu do publikacji nie mogą naruszać zasadniczej myśli autorów. Wszelkie modyfikacje o charakterze merytorycznym są z nimi konsultowane.
6. W przypadku podjęcia decyzji o niepublikowaniu artykułu nie może on zostać w żaden sposób wykorzystany przez wydawcę lub uczestników procesu wydawniczego bez pisemnej zgody autorów. Autorzy mogą się odwołać od decyzji o niepublikowaniu artykułu. W tym celu powinni się skontaktować z redaktorem naczelnym lub sekretarzem redakcji „WS” i przedstawić stosowną argumentację. Odwołania autorów są rozpatrywane przez redaktora naczelnego.
7. Członkowie RN i KR ani pracownicy WCN nie mogą pozostawać w jakimkolwiek konflikcie interesów w odniesieniu do artykułów zgłaszanych do publikacji. Przez konflikt interesów należy rozumieć sytuację, w której jakiekolwiek interesy lub zależności (służbowe, finansowe lub inne) mogą mieć wpływ na ocenę artykułu lub decyzję o jego publikacji.
8. W celu przeciwdziałania nierzetelności naukowej wymagane jest złożenie przez autorów oświadczenia, w którym deklarują, że zgłaszany artykuł nie narusza praw autorskich osób trzecich, nie był dotychczas publikowany i jest ich oryginalnym dziełem, a także określają swój wkład w opracowanie artykułu.
9. W celu zapewnienia wysokiej jakości recenzji wymagane jest złożenie przez recenzentów oświadczenia o przestrzeganiu zasad etyki recenzowania COPE i niewystępowaniu konfliktu interesów.
10. W przypadku uzasadnionego podejrzenia na jakimkolwiek etapie procesu wydawniczego, że autorzy dopuścili się nierzetelności naukowej (zob. pkt 3.1. Odpowiedzialność autorów), zespół redakcyjny skrupulatnie zbada sprawę ewentualnego nadużycia. Jeśli nierzetelność autorów zostanie udowodniona, to zgłoszony przez nich artykuł zostanie odrzucony przez KR, a autorzy otrzymają informację o podjętej decyzji wraz z jej uzasadnieniem.
11. Czytelnicy, którzy mają wobec autorów opublikowanego artykułu uzasadnione podejrzenia o nierzetelność naukową, powinni powiadomić o tym redaktora naczelnego lub sekretarza redakcji. Po zbadaniu sprawy ewentualnego nadużycia czytelnicy zostaną poinformowani o rezultacie przeprowadzonego postępowania. W przypadku potwierdzenia nadużycia, na łamach czasopisma zostanie zamieszczona stosowna informacja.

3.3. Odpowiedzialność recenzentów

1. Recenzenci przyjmują artykuł do recenzji tylko wtedy, gdy uznają, że:
 - posiadają odpowiednią wiedzę w określonej dziedzinie, aby rzetelnie ocenić pracę;
 - zgodnie z ich stanem wiedzy nie istnieje konflikt interesów w odniesieniu do autorów, przedstawionych w artykule badań i instytucji je finansujących, co potwierdzają w oświadczeniu;
 - mogą wywiązać się z terminu ustalonego przez redakcję, aby nie opóźnić publikacji.

2. Recenzenci są zobligowani do zachowania obiektywności i poufności oraz powstrzymania się od osobistej krytyki. Zawsze powinni uzasadnić swoją ocenę, przedstawiając stosowną argumentację.
3. Recenzenci powinni wskazać ważne dla wyników badań opublikowane prace, które w ich ocenie powinny zostać przywołane w ocenianym artykule.
4. W razie stwierdzenia wysokiego poziomu zbieżności treści recenzowanej pracy z innymi opublikowanymi materiałami lub podejrzenia innych przejawów nierzetelności naukowej recenzenci są zobowiązani poinformować o tym redakcję.
5. Po ukończeniu recenzji przechowywanie przesłanych przez redakcję materiałów (w jakiegokolwiek formie) oraz posługiwanie się nimi przez recenzentów jest niedozwolone.

3.4. Odpowiedzialność wydawcy

1. Materiały opublikowane w „WS” są chronione prawem autorskim.
2. Wydawca udostępnia pełną treść wszystkich artykułów w Internecie w trybie otwartego dostępu, tj. bezpłatnie i bez technicznych ograniczeń. Użytkownicy mogą czytać, pobierać, kopiować, drukować i wykorzystywać do innych celów artykuły zamieszczone online, zgodnie z właściwymi przepisami o dozwolonym użytku, pod warunkiem wskazania źródła pochodzenia artykułu. Inne sposoby wykorzystania treści artykułów „WS” wymagają zgody wydawcy.
3. Wydawca deklaruje gotowość do opublikowania poprawek, wyjaśnień oraz przeprosin.

4. Wymogi redakcyjne

Zgodnie z wymogami czasopisma omawiany w artykule problem badawczy powinien być jednoznacznie zdefiniowany oraz istotny dla oceny zjawisk społecznych lub gospodarczych. Artykuł powinien zawierać wyraźnie określony cel badania, precyzyjny opis badanych zjawisk i stosowanych metod, uzyskane wyniki przeprowadzonej analizy oraz autorskie wnioski.

4.1. Struktura i zawartość artykułu

Wymagane elementy artykułu:

1. Tytuł.
2. Dane autora: imię i nazwisko, ORCID, adres e-mail, afiliacja. Wśród autorów artykułu wieloautorskiego należy wskazać autora korespondencyjnego.
3. Streszczenie (zalecana objętość – do 1200 znaków ze spacjami, forma bezosobowa). W przypadku artykułu opisującego badanie empiryczne powinno zawierać: cel, przedmiot, okres i metodę badania, źródła danych i najważniejsze wnioski z badania. W przypadku artykułów o innym charakterze należy podać co najmniej cel pracy, jej przedmiot i najważniejsze wnioski.

Streszczenie to podstawowe źródło informacji o artykule, warunkujące też decyzję czytelnika o zapoznaniu się z całą pracą. Dlatego powinno być przygotowane ze szczególną starannością i dbałością o umieszczenie w nim wszystkich wymaganych elementów.

4. Słowa kluczowe – najistotniejsze pojęcia lub wyrażenia użyte w pracy (nie mniej niż trzy). Powinny być zawarte w streszczeniu i/lub tytule.
5. Kod/kody z klasyfikacji Journal of Economic Literature (JEL).
6. Tłumaczenie tytułu, streszczenia i słów kluczowych (na język angielski w przypadku artykułu napisanego w języku polskim, a na język polski w przypadku artykułu napisanego w języku angielskim).
7. W artykule opisującym badanie empiryczne wymagane są następujące części:
 - wprowadzenie, zawierające: syntetyczne przedstawienie zagadnień teoretycznych, uzasadnienie podjęcia danego problemu badawczego, cel badania i krytyczne odniesienie do literatury przedmiotu. W wyjątkowych przypadkach, kiedy istotne dla podjętego tematu jest obszerniejsze przedstawienie dyskusji toczącej się w literaturze, przegląd literatury może stanowić odrębną część artykułu;
 - metoda badania, zawierająca: przedmiot i okres badania, źródła danych i zastosowane metody badawcze, w tym uzasadnienie ich wyboru;
 - wyniki badania wraz z wnioskami;
 - podsumowanie, które powinno być zwięzłe i odzwierciedlać istotę problemu badawczego przedstawionego w artykule, bez podawania danych liczbowych; końcowe wnioski powinny odnosić się do treści artykułu, a w szczególności do celu badania.Wszystkie części powinny być opatrzone numerami.
8. Bibliografia, zawierająca pełny wykaz prac i materiałów przywołanych w artykule, przygotowana zgodnie z wymogami czasopisma.

4.2. Przygotowanie artykułu

1. Artykuł powinien być utrzymany w formie bezosobowej.
2. Tekst należy zapisać alfabetem łacińskim. Nazwy własne, tytuły itp. oryginalnie zapisane innym alfabetem powinny być poddane transliteracji.
3. Nie należy stosować stylów; formatowanie należy ograniczyć do wymogów redakcyjnych.
4. Objętość artykułu łącznie ze streszczeniem, słowami kluczowymi, bibliografią, tablicami, wykresami i innymi materiałami graficznymi nie powinna być mniejsza niż 10 stron maszynopisu ani przekraczać 20 stron.
5. Edytor tekstu: Microsoft Word, format *.doc lub *.docx.
6. Krój czcionki:
 - Arial – tytuł, autor, streszczenia, słowa kluczowe, kody JEL, śródtytuły, elementy graficzne (tablice, zestawienia, wykresy, schematy), przypisy;
 - Times New Roman – tekst główny, bibliografia.
7. Wielkość czcionki:
 - 14 pkt – tytuł, autor, tytuły rozdziałów;
 - 12 pkt – tekst główny, tytuły podrozdziałów;
 - 10 pkt – pozostałe elementy.
8. Marginesy – 2,5 cm z każdej strony.
9. Interlinia – 1,5 wiersza; tablice i przypisy – 1 wiersz; przed tytułami rozdziałów i podrozdziałów oraz po nich – pusty wiersz.
10. Wcięcie akapitowe – 0,4 cm; bibliografia – bez wcięcia, wysunięcie 0,4 cm.

11. Przy wyliczeniach należy posłużyć się listą punktowaną z punktorami w postaci kropek (wysunięcie 0,4 cm, wcięcie 0 cm); wiersze (oprócz ostatniego) zakończone średnikiem.
12. Strony ponumerowane automatycznie.
13. Tablice i elementy graficzne (wykresy, mapy, schematy) muszą być przywołane w tekście.
14. Wykresy, mapy i schematy należy zamieścić w tekście głównym. Wykresy powinny być edytowalne (optymalnie wykonane w programie Excel; w przypadku wykonania w programie graficznym powinny mieć postać wektorową). Dane, na podstawie których opracowano wykresy, należy przekazać osobno w pliku programu Excel (lub innym edytowalnym w pakiecie Microsoft Office), ewentualnie wykresy powinny dawać możliwość odczytania z nich danych.
15. Tablice muszą być edytowalne. Nie należy stosować rastrów, cieniowania, pogrubiania czy też podwójnych linii itp.
16. Wskazówki dotyczące opracowywania map znajdują się w publikacji *Mapy statystyczne. Opracowanie i prezentacja danych*, dostępnej na stronie internetowej GUS.
17. Pod tablicami i każdym elementem graficznym należy podać źródło opracowania, a także objaśnić użyte w nich skróty i symbole.
18. Literowe symbole liczb i innych wielkości niezłożonych należy zapisywać małą lub dużą literą i pismem pochyłym (np. a , A , $y(x)$, a_i); wektorów – małą literą, pismem pochyłym i pogrubionym (np. $\mathbf{a}$, $\mathbf{w}$, $\mathbf{y}(x)$, $\mathbf{w}_i$); macierzy – dużą literą i pismem pogrubionym (np. $\mathbf{A}$, $\mathbf{M}$, $\mathbf{Y}(x)$, $\mathbf{M}_i$).
19. Objaśnienia znaków umownych w tablicach: kreska (–) – zjawisko nie wystąpiło; zero (0) – zjawisko istniało w wielkości mniejszej od 0,5; (0,0) – zjawisko istniało w wielkości mniejszej od 0,05; kropka (.) – brak informacji, konieczność zachowania tajemnicy statystycznej, wypełnienie pozycji jest niemożliwe lub niecelowe; „w tym” – oznacza, że nie podaje się wszystkich składników sumy.
20. Stosowane są następujące skróty: tablica – tabl., wykres – wykrr.
21. Wszystkie zawarte w artykule informacje, dane i stwierdzenia wykraczające poza wiedzę powszechną – np. wyniki badań innych autorów, zarówno o charakterze empirycznym, jak i koncepcyjnym – muszą być opatrzone przypisem bibliograficznym. Przez wiedzę powszechną należy rozumieć informacje ogólnie znane i niebudzące wątpliwości ani kontrowersji w danej grupie społecznej, np. utworzenie GUS w 1918 r. lub powstanie UE w 1993 r. na podstawie traktatu z Maastricht. Natomiast dane statystyczne udostępniane lub publikowane np. przez GUS lub Eurostat nie należą do takich informacji. Charakteru wiedzy powszechnej nie mają również stwierdzenia odnoszące się do idei, zjawisk i procesów społecznych, politycznych czy gospodarczych. Nawet pozornie zdroworozsądkowe idee zmieniają bowiem swój sens w zależności od kultury, języka lub dyscypliny naukowej, a także bywają w rozmaity sposób konceptualizowane, jak np. pojęcie poznania w naukach społecznych.

Podanie źródła jest konieczne niezależnie od tego, czy informacje lub stwierdzenia są ujęte w ramy cytatu, czy przedstawione bez dosłownego przytoczenia, np. w formie parafrazy. Jeżeli stwierdzenie może budzić jakiegokolwiek wątpliwości odbiorców, autor powinien wskazać stosowne źródło podawanej informacji.

22. Przypisy rzeczowe, słownikowe lub informacyjne należy umieszczać na dole strony. Przypisy bibliograficzne, zgodnie ze standardem APA (American Psychological Association), należy podawać w tekście głównym.
23. Bibliografię należy przygotować zgodnie ze standardem APA.

4.3. Zasady przywoływania publikacji w treści artykułu

Wyszczególnienie	Przykład przywołania	
	w odsyłaczu	w treści zdania
Autor indywidualny		
Jeden autor	(Iksiński, 2001)	Iksiński (2001)
Dwóch autorów	(Iksiński i Nowak, 1999)	Iksiński i Nowak (1999)
Trzech autorów lub więcej	(Jankiewicz i in., 2003)	Jankiewicz i in. (2003)
Autor instytucjonalny		
Nazwa funkcjonuje jako powszechnie znany skrótowiec: pierwsze przywołanie w tekście	(International Labour Organization [ILO], 2020)	International Labour Organization (ILO, 2020)
kolejne przywołanie	(ILO, 2020)	ILO (2020)
Pełna nazwa	(Stanford University, 1995)	Stanford University (1995)
Typ publikacji		
Publikacja bez ustalonego autorstwa	(Skrócony tytuł ..., 2015)	Pełny tytuł (2015)
Publikacja bez roku wydania	(Iksiński, b.r.)	Iksiński (b.r.)
Akt prawny	(Pełny tytuł)	Pełny tytuł
Strona internetowa / Zbiór danych: znana data publikacji	(Iksiński, 2020) / (Nazwa instytucji, 2020)	Iksiński (2020) / Nazwa instytucji (2020)
nieznana data publikacji	(Iksiński, b.r.) / (Nazwa instytucji, b.r.)	Iksiński (b.r.) / Nazwa instytucji (b.r.)
Rodzaj przywołania		
Przywoływanie kilku prac (porządek prac – chronologiczny, porządek autorów – alfabetyczny)	(Iksiński, 1997, 1999, 2004a, 2004b; Nowak, 2002)	Iksiński (1997, 1999, 2004a, 2004b) i Nowak (2002)
Przywoływanie publikacji za innym autorem (uwaga: w bibliografii należy wymienić tylko pracę cytowaną)	(Nowakowski, 1990, za: Zieniecka, 2007)	Nowakowski (1990, za: Zieniecka, 2007)

Źródło: opracowanie na podstawie: American Psychological Association. (2020). *Publication manual of the American Psychological Association* (7th edition). <https://doi.org/10.1037/0000165-000>.

4.4. Przykłady opisu bibliograficznego

Bibliografia powinna być zamieszczona na końcu opracowania. Prace należy podać alfabetycznie według nazwiska pierwszego autora. W przypadku dwóch lub więcej prac tego samego autora / tych samych autorów trzeba je uporządkować chronologicznie według roku publikacji. Jeśli kilka prac tego samego autora / tych samych autorów zostało opublikowanych w tym samym roku, należy ułożyć je alfabetycznie według tytułu i odpowiednio oznaczyć literami a, b, c itd.

Typ publikacji	Przykład opisu bibliograficznego
Artykuł w czasopiśmie	
W wersji drukowanej	Nazwisko, X. (rok). Tytuł artykułu. <i>Tytuł czasopisma, rocznik</i> (zeszyt), strona początku–strona końca.
Dostępny w internecie, z DOI	Nazwisko, X., Nazwisko 2, Y. (rok). Tytuł artykułu. <i>Tytuł czasopisma, rocznik</i> (zeszyt), strona początku–strona końca. https://doi.org/xxx .
Dostępny w internecie, bez DOI	Nazwisko, X., Nazwisko 2, Y., Nazwisko 3, Z. (rok). Tytuł artykułu. <i>Tytuł czasopisma, rocznik</i> (zeszyt), strona początku–strona końca. https://xxx .
Maszynopis	
Niepublikowany / w przygotowaniu / zgłoszony do publikacji	Nazwisko, X. (rok). <i>Tytuł [maszynopis niepublikowany / w przygotowaniu / zgłoszony do publikacji]</i> . Nazwa instytucji, w której powstaje lub powstał maszynopis.
Opublikowany nieformalnie	Nazwisko, X., Nazwisko 2, Y. (rok). <i>Tytuł artykułu</i> . https://xxx .
Książka	
W wersji drukowanej	Nazwisko, X. (rok). <i>Tytuł książki</i> . Miejsce wydania: Wydawnictwo.
Dostępna w internecie, z DOI	Nazwisko, X. (rok). <i>Tytuł książki</i> . Miejsce wydania: Wydawnictwo. https://doi.org/xxx .
Dostępna w internecie, bez DOI	Nazwisko, X. (rok). <i>Tytuł książki</i> . Miejsce wydania: Wydawnictwo. https://xxx .
W przekładzie	Nazwisko, X. (rok). <i>Tytuł książki</i> (tłum. Y. Nazwisko). Miejsce wydania: Wydawnictwo.
Wydanie wielotomowe: tom zatytułowany	Nazwisko, X. (rok). <i>Tytuł książki: nr tomu. Tytuł tomu</i> . Miejsce wydania: Wydawnictwo.
tom niezatytułowany	Nazwisko, X. (rok). <i>Tytuł książki</i> (nr tomu). Miejsce wydania: Wydawnictwo.
Kolejne wydanie	Nazwisko, X. (rok). <i>Tytuł książki</i> (nr wydania). Miejsce wydania: Wydawnictwo.
Pod redakcją: w języku polskim	Nazwisko, X. (red.). (rok). <i>Tytuł książki</i> . Miejsce wydania: Wydawnictwo.
w języku angielskim	Nazwisko, X. (Ed.). (rok). <i>Tytuł książki</i> . Miejsce wydania: Wydawnictwo.
Rozdział w pracy zbiorowej	Nazwisko, X. (rok). Tytuł rozdziału. W: Y. Nazwisko, Z. Nazwisko 2 (red.), <i>Tytuł książki</i> (s. strona początku–strona końca). Miejsce wydania: Wydawnictwo. https://doi.org/xxx lub https://xxx .
Inne prace	
Raport: autor indywidualny	Nazwisko, X. (rok). <i>Tytuł raportu</i> . Miejsce wydania: Wydawnictwo.
autor instytucjonalny	Nazwa instytucji. (rok). <i>Tytuł raportu</i> . Miejsce wydania: Wydawnictwo.
Working Papers	Nazwisko, X. (rok). <i>Tytuł pracy</i> (nazwa serii i numer). https://doi.org/xxx lub https://xxx .
Materiały z konferencji: nieopublikowane	Nazwisko. X. (rok, dzień i miesiąc). <i>Tytuł pracy</i> [typ wystąpienia, np. referat]. Nazwa konferencji, miejsce konferencji.
opublikowane	Nazwisko. X. (rok). <i>Tytuł pracy</i> . Nazwa konferencji, miejsce konferencji.
Rozprawa doktorska: nieopublikowana	Nazwisko, X. (rok). <i>Tytuł pracy</i> [niepublikowana rozprawa doktorska]. Nazwa instytucji nadającej tytuł doktorski.
opublikowana	Nazwisko, X. (rok). <i>Tytuł pracy</i> [rozprawa doktorska, nazwa instytucji nadającej tytuł doktorski]. https://xxx .
Akt prawny	Pełny tytuł aktu prawnego wraz z datą publikacji w dzienniku urzędowym.

Typ publikacji	Przykład opisu bibliograficznego
Strona internetowa	
Znana data publikacji, zawartość strony się nie zmienia	Nazwisko, X. (rok, dzień i miesiąc). <i>Tytuł</i> . https://xxx .
Nieznana data publikacji, zawartość strony się zmienia	Nazwa instytucji. (b.r.). <i>Tytuł</i> . Pobrane dzień, miesiąc i rok pobrania z https://xxx .
Zbiór danych	
Surowe dane nieopublikowane	Nazwisko, X. (rok wydania pracy, w której dane są wykorzystywane) [opis danych, np. surowe dane nieopublikowane dotyczące...]. Źródło danych (np. nazwa uniwersytetu).
Dane opublikowane: znana data publikacji, zawartość zbioru się nie zmienia	Nazwisko, X. (rok). <i>Nazwa zbioru danych</i> [zbiór danych]. Wydawca. https://xxx .
nieznana data publikacji, zawartość zbioru się zmienia	Nazwa instytucji. (b.r.). <i>Nazwa zbioru danych</i> [zbiór danych]. Wydawca. Pobrane dzień, miesiąc i rok pobrania z https://xxx .

Źródło: opracowanie na podstawie: American Psychological Association. (2020). *Publication manual of the American Psychological Association* (7th edition). <https://doi.org/10.1037/0000165-000>.

Praca przygotowana w sposób niezgodny z powyższymi wskazówkami będzie odesłana do Autora z prośbą o dostosowanie formy artykułu do wymogów redakcyjnych.

ZAKRES TEMATYCZNY DZIAŁÓW

THEMATIC SCOPE OF SECTIONS

(for the English translation of the information given below, please visit ws.stat.gov.pl/AimScope)

Studia metodologiczne

W tym dziale zamieszczane są artykuły naukowe przedstawiające teoretyczne rozwiązania metodologiczne ze wskazaniem ich praktycznej użyteczności, w tym prace przeglądowe i porównawcze oraz dotyczące etyki w statystyce. Poruszane w nich zagadnienia obejmują różne dziedziny statystyki, ekonomii matematycznej i ekonometrii. Omawiane rezultaty badawcze mogą znaleźć efektywne zastosowanie w badaniach empirycznych oraz analizach statystycznych i służyć podnoszeniu ich jakości, jak również powiększeniu zasobu informacyjnego.

Statystyka w praktyce

Dział ten zawiera artykuły poświęcone nowatorskim zastosowaniom w praktyce znanych narzędzi i modeli statystycznych oraz analizie i ocenie statystycznej zjawisk społeczno-ekonomicznych i innych; zamieszczane tu prace opierają się w szczególności na danych pochodzących z zasobów statystyki publicznej. Zastosowania w praktyce obejmują również wykorzystanie narzędzi informatycznych do uzyskiwania i przetwarzania informacji statystycznych, naliczania danych wynikowych, ich prezentacji i rozpowszechniania. Może to też dotyczyć opracowań stosujących nowoczesne techniki programistyczne pozwalające na efektywną komunikację z systemami informacyjnymi oraz ułatwiające wykorzystanie danych wynikowych. Publikowane są także artykuły sygnalizujące problemy związane z projektowaniem badań statystycznych, uzyskiwaniem, integracją i przetwarzaniem danych oraz generowaniem wynikowych informacji statystycznych i kontrolą ich ujawniania wraz z propozycjami efektywnych rozwiązań w tym zakresie.

Studia interdyscyplinarne. Wyzwania badawcze

To blok tematyczny zawierający artykuły wskazujące i podejmujące wyzwania badawcze, które są szczególnie istotne ze względu na rosnące potrzeby współczesnych użytkowników danych statystycznych i wymagają zaangażowania znacznych nakładów pracy, środków oraz rozwiązań z różnych dziedzin nauki i techniki. W dziale tym publikowane są również opracowania dotyczące: wykorzystania technologii informacyjnych i komunikacyjnych (ICT), gospodarki opartej na wiedzy, problematyki innowacyjności, przepływu informacji we współczesnym społeczeństwie oraz przetwarzania i analizy zagadnień związanych z data science i big data, a zatem problematyki bardzo często powiązanej z działaniami interdyscyplinarnymi.

Edukacja statystyczna

W tym dziale zamieszczane są artykuły dotyczące metod i efektów nauczania statystyki oraz popularyzacji myślenia statystycznego. Odnosi się to zwłaszcza do problemów związanych z kształceniem w zakresie umiejętności stosowania statystyki na wszystkich poziomach edukacji, a także do wykorzystywania nowoczesnych koncepcji i metod dydaktycznych oraz pomocy naukowych w nauczaniu statystyki. Uwaga skoncentrowana jest na rozumieniu prawdopodobieństwa i statystyki, badaniach z zakresu nauczania statystyki, postaw i zachowań społecznych w odniesieniu do tej dziedziny wiedzy, jak również na rozumieniu informacji statystycznych. Ponadto ukazywane są problemy związane z prezentacją danych statystycznych oraz ich interpretacją w powszechnym obiegu informacyjnym, np. w środkach społecznego przekazu.

Z dziejów statystyki

Prace publikowane w tym dziale poświęcone są historii prowadzenia obserwacji statystycznych oraz rozwoju ich metodologii i narzędzi. Ponadto zamieszczane są tu informacje dotyczące życia i osiągnięć zawodowych wybitnych statystyków, jak również najważniejszych instytucji i organizacji statystycznych w Polsce i za granicą.

Dyskusje. Recenzje. Informacje

Jedyny dział zawierający teksty nierecenzowane i niemające charakteru artykułów naukowych. Obejmuje informacje o najważniejszych wydarzeniach dotyczących statystyki polskiej i międzynarodowej, a także sprawozdania z konferencji naukowych, recenzje książek i opracowań z zakresu statystyki i jej zastosowań, rekomendacje nowych, istotnych i ciekawych pozycji wydawniczych z tego obszaru wiedzy, jak również odpowiedzi autorów na recenzje oraz polemiki, dyskusje i sprostowania dotyczące artykułów zamieszczonych na łamach czasopisma.